

Université Mohammed V - Agdal
Faculté des Sciences
Département de Mathématiques et Informatique
Avenue Ibn Batouta, B.P. 1014
Rabat, Maroc

.:: Module Mathématiques I : Algèbre ::.

Filière :

Sciences de Matière Physique (SMP)

et

Sciences de Matière Chimie(SMC)

Chapitre I: Opérations logiques élémentaires. Ensembles.
Quantificateurs. Relations binaires et Applications

Par

Prof: Jilali Mikram
Groupe d'Analyse Numérique et Optimisation
<http://www.fsr.ac.ma/ANO/>
Email : mikram@fsr.ac.ma

Année : 2005-2006

TABLE DES MATIERES

1 Opérations logiques élémentaires. Ensembles. Quantificateurs. Relations binaires et Applications	3
1.1 Qu'est-ce qu'une expression mathématique ?	3
1.2 Négation d'une proposition : non P	3
1.3 Disjonction de deux propositions : P ou Q	4
1.4 Equivalence de deux propositions : $P \Leftrightarrow Q$	5
1.5 Conjonction de deux propositions : P et Q	5
1.6 Implication logique de deux propositions : $P \Rightarrow Q$	6
2 Quelques formes de raisonnements	7
2.1 Raisonnement à partir de la contraposée	7
2.2 Raisonnement par l'absurde	8
3 Notions sur les ensembles	8
3.1 Définition d'un ensemble	8
3.2 Définition d'un sous-ensemble et de l'ensemble vide	9
3.3 Intersection et union d'ensembles	10
3.4 Complémentaire d'une partie d'un ensemble	10
3.5 Cardinal d'un ensemble	11
3.6 Produit cartésien d'ensembles	12
4 Quantificateurs	13
4.1 Proposition dépendant d'une variable : $P(x)$	13
4.2 Quantificateur universel : $\forall$	13
4.3 Quantificateur existentiel : $\exists$	14
4.4 Quantificateurs multiples	14
5 Relation binaires	15
5.1 Relation d'équivalence	16
5.2 Classes d'équivalence	16
5.3 Relations d'ordre	17
5.4 Eléments particuliers d'un ensemble ordonné	18
6 Quelques mots sur les applications	19
6.1 Composition des applications	21
6.2 Définition de l'application réciproque	22
6.3 Composition des applications réciproques	23

1 Opérations logiques élémentaires. Ensembles. Quantificateurs. Relations binaires et Applications

1.1 Qu'est-ce qu'une expression mathématique ?

Il s'agit de se familiariser à l'expression mathématique du raisonnement et à quelques règles de raisonnement logique constamment utilisées en mathématiques et ailleurs.

Ces règles permettent, à partir de propositions sur (ou propriétés, ou relations entre) des objets mathématiques (nombres, figures géométriques, fonctions, . . .), connues ou posées comme vraies, de déduire d'autres propositions ou propriétés vraies.

Ici, le mot 'proposition' désigne tout énoncé sur les objets considérés auquel on peut attribuer une valeur de vérité. Par exemple :

- (P_1) $\sqrt{2}$ est un nombre rationnel,
- (P_2) par deux points il passe une droite et une seule,
- (P_3) une fonction dérivable est continue.

Quant à la vérité en question, il s'agit d'une valeur logique qui est l'un des deux mots vraie ou fausse (principe du tiers exclu). Ainsi (P_1) est fausse et (P_3) est vraie.

Un certain nombre de propositions sont considérées comme vérités premières, qui ne se déduisent pas d'autres propositions vraies. Certaines d'entre elles traduisent en langage mathématique les propriétés les plus évidentes des objets concrets auxquels on pense. On les appelle des axiomes. Par exemple, (P_2) est un des axiomes de la géométrie euclidienne.

Les autres propositions vraies le sont par déduction des axiomes ou d'autres propositions dont la vérité est déjà démontrée. Les axiomes sont en petit nombre et possèdent une cohérence interne, en ce sens qu'on ne peut déduire d'eux aucune proposition à la fois vraie et fausse.

1.2 Négation d'une proposition : non P

Définition 1.1 Si P est une proposition, sa négation, notée $\text{non}(P)$ est une proposition qui est fausse si P est vraie et qui est vraie si P est fausse.

Il résulte de cette définition que $\text{non}(\text{non}P)$ et P ont la même valeur logique, c'est à dire sont vraies simultanément ou fausses simultanément.

Par exemple:

- P : Tous les dimanches je vais au restaurant,

$\text{non}(P)$: Il existe au moins un dimanche où je ne vais pas au restaurant
 – P : Je vais au restaurant au moins un dimanche de décembre,
 $\text{non}(P)$: Je ne vais jamais au restaurant le dimanche en décembre.
 Remarque I.1.1. $\text{non}(P)$ se note aussi $\overline{P}$.

1.3 Disjonction de deux propositions : P ou Q

Définition 1.2 Si P et Q sont deux propositions, la disjonction, notée P ou Q , est une proposition qui est vraie si au moins l'une des deux propositions est vraie et qui est fausse si les deux propositions sont fausses.

On introduit maintenant la notion de table de vérité :

Définition 1.3 Soient P et Q deux propositions, R une proposition dépendant de P et Q (dans cet ordre). On associe à R le tableau suivant :

	V	F
V		
F		

où les cases blanches seront remplies par la lettre V chaque fois que R est vraie et la lettre F si elle est fausse, selon les valeurs logiques V ou F de P et Q respectivement indiquées sur la 1^{ère} colonne et la 1^{ère} ligne. Ce tableau s'appelle la table de vérité de R .

Par exemple si $R = (P \text{ ou } Q)$, sa table de vérité s'écrit :

	V	F
V	V	V
F	V	F

Par exemple, si on considère les deux propositions :

- P : Tous les lundis je vais au cinéma,
- Q : Le 15 de chaque mois je vais au cinéma,

La proposition (P ou Q) est vraie si elle s'applique à quelqu'un qui va au cinéma tous les lundis ou à quelqu'un qui va au cinéma le 15 de chaque mois (il peut évidemment faire les deux). Elle est fausse dans tous les autres cas. En particulier elle est fausse s'il s'agit de quelqu'un qui ne va au cinéma que les lundis 15.

Attention, le 'fromage ou dessert' du restaurant n'est pas un 'ou mathématique' car il est exclusif.

Si dans une démonstration on veut utiliser l'hypothèse P ou Q est vraie, alors deux cas sont possibles :

- soit P est vraie et on utilise ce résultat dans la démonstration,
- soit P est fausse, alors Q est vraie et l'on utilise ces deux résultats dans la démonstration.

Pour montrer que P ou Q est vraie, il faut démontrer que l'on est dans l'un des deux cas suivants :

- soit P est vraie et donc P ou Q est vraie,
- soit P est fausse et ceci peut être utilisé pour montrer que Q est vraie.

Remarque 1.1 P ou Q se note aussi $P \vee Q$

1.4 Equivalence de deux propositions : $P \Leftrightarrow Q$

Définition 1.4 Deux propositions P et Q sont équivalentes si elles ont la même valeur logique. On note $P \Leftrightarrow Q$

Lorsque P et Q sont équivalentes, on dit aussi que P (resp Q) est une condition nécessaire et suffisante de Q (resp P), ou que P (resp. Q) est vraie si et seulement si Q (resp. P) est vraie.

Dans ce cas, les deux propositions sont vraies ou fausses simultanément. Par exemple :

$$\begin{aligned}\{\text{non (non } P)\} &\Leftrightarrow P \\ \{P \text{ ou } Q\} &\Leftrightarrow \{Q \text{ ou } P\} \\ \{2x = 4\} &\Leftrightarrow \{x = 2\}\end{aligned}$$

La disjonction est associative dans le sens où

$$\{P \text{ ou } (Q \text{ ou } R)\} \Leftrightarrow \{(P \text{ ou } Q) \text{ ou } R\}$$

Si P est toujours fausse, alors $(P \text{ ou } Q)$ est équivalente à Q .

1.5 Conjonction de deux propositions : P et Q .

Définition 1.5 Si P et Q sont deux propositions, la conjonction, notée $(P \text{ et } Q)$, est la proposition qui est vraie si les deux propositions sont vraies et qui est fausse si au moins l'une des deux propositions est fausse.

Il résulte de cette définition que les propositions $(P \text{ et } Q)$ et $(Q \text{ et } P)$ sont équivalentes.

Par exemple, soient les deux propositions :

– P : Tous les lundis je vais au cinéma,

– Q : Le 15 de chaque mois je vais au cinéma,

La proposition $(P \text{ et } Q)$ est vraie si elle s'applique à quelqu'un qui va au cinéma tous les lundis et le 15 de chaque mois. Elle est fausse dans tous les autres cas. Attention, $(P \text{ et } Q)$ ne correspond pas à : Tous les lundis 15 je vais au cinéma.

La conjonction est associative dans le sens où $((P \text{ et } Q) \text{ et } R)$ est équivalente à $(P \text{ et } (Q \text{ et } R))$. Si P est toujours vraie, alors $(P \text{ et } Q)$ est équivalente à Q .

Remarque I.1.3. $(P \text{ et } Q)$ se note aussi $P \wedge Q$.

1.6 Implication logique de deux propositions : $P \Rightarrow Q$

Définition 1.6 Soient P et Q deux propositions, on appelle l'implication logique (de Q par P) la proposition, notée $P \Rightarrow Q$, qui est vraie si

– soit P est fausse,

– soit P est vraie et Q est vraie.

Elle est fausse dans le seul cas où P est vraie et Q est fausse.

Attention $P \Rightarrow Q$ n'est pas équivalente à $Q \Rightarrow P$. L'implication se dit, en langage courant, ' P implique Q ' et signifie que Q est vraie dès que P l'est. D'ailleurs pour prouver que cette implication est vraie, on n'a qu'une seule chose à faire : démontrer que si P est vraie, alors Q aussi l'est. Mais il faut faire attention car elle ne donne aucun renseignement sur Q si P est fausse, comme on le voit dans l'exemple suivant :

Soient 3 nombres réels x, y, z . On a l'implication (bien connue) suivante :

$$(x = y) \Rightarrow (xz = yz)$$

On voit sur cet exemple que quand la proposition (P) est fausse ($x \neq y$) la conclusion (Q) peut être vraie (si $z = 0$) ou fausse (si $z \neq 0$) .

Dans la pratique, par abus de langage, quand la notation $(P \Rightarrow Q)$ est utilisée, on entend que cette implication est vraie : on dira 'démontrer $P \Rightarrow Q$ ' plutôt que 'démontrer que $(P \Rightarrow Q)$ est vraie.

Proposition 1.1 $(P \Rightarrow Q)$ est équivalente à $((\text{non } P) \text{ ou } Q)$.

Démonstration: On a vu que $((\text{non } P) \text{ ou } Q)$ est vraie si $(\text{non } P)$ est vraie (donc P fausse) ou si $(\text{non } P)$ est fausse (P vraie) et Q est vraie. Ceci correspond bien à $(P \Rightarrow Q)$

Au lieu de dire que P implique Q on dit aussi que P est une condition suffisante de Q (pour que Q soit vraie, il suffit que P le soit), ou que Q est une condition nécessaire de P (si P est vraie, nécessairement Q l'est).

Corollaire 1.1 *non $(P \Rightarrow Q)$ est équivalente à $(P$ et $(\text{non } Q))$.*

Attention ! La négation d'une implication n'est pas une implication.

Enfin, l'implication est transitive, soit

$$((P_1 \Rightarrow P_2) \text{ et } (P_2 \Rightarrow P_3)) \Rightarrow (P_1 \Rightarrow P_3)$$

Proposition 1.2 $(P \Rightarrow Q) \Leftrightarrow (\text{non } Q) \Rightarrow \text{non } P$

La démonstration est à faire en exercice.

Proposition 1.3 $(P \Leftrightarrow Q)$ est équivalent à $P \Rightarrow Q$ et $Q \Rightarrow P$.

L'équivalence est transitive, soit:

$$((P_1 \Leftrightarrow P_2) \text{ et } (P_2 \Leftrightarrow P_3)) \Rightarrow (P_1 \Leftrightarrow P_3)$$

2 Quelques formes de raisonnements

Un raisonnement est une manière d'arriver à une conclusion en partant d'une (ou de plusieurs) hypothèse(s), et en utilisant les règles de déduction d'une proposition à partir d'une autre. Vous connaissez déjà le raisonnement par équivalence qui consiste à partir d'une proposition vraie (l'hypothèse par exemple) et à construire par équivalence d'autres propositions (qui sont donc vraies), dont la dernière est la conclusion. Vous connaissez le raisonnement par récurrence. Voici deux autres formes de raisonnement qui découlent des règles de logique précédentes.

2.1 Raisonnement à partir de la contraposée

La proposition 1.2 donne une autre manière de démontrer que $P \Rightarrow Q$. En effet il est équivalent de montrer que $(\text{non } Q) \Rightarrow (\text{non } P)$, c'est-à-dire que si la proposition Q est fausse alors la proposition P est fausse, ce qui est parfois plus simple. C'est ce que l'on appelle un raisonnement par contraposée. Un premier exemple emprunté à Racine est : 'Si Titus est jaloux, il est amoureux'. En effet, s'il n'est pas amoureux, il n'a aucune raison d'être jaloux !

Un deuxième exemple mathématique : Si n est un entier impair alors le chiffre des unités de n est impair. On va montrer la contraposée, à savoir : (le chiffre des unités de n est pair) $\Rightarrow n$ est pair.

En effet, si le chiffre des unités de n est pair, on peut écrire $n = 10q + 2p$ soit $n = 2(5q + p)$ c'est-à-dire n est pair.

Attention La proposition $(P \Rightarrow Q) \Leftrightarrow (nonP \Rightarrow nonQ)$ est fausse!. Elle peut conduire à de nombreuses erreurs, par exemple la suivante : étant donné que 'tout homme est mortel', cet énoncé pourrait servir à prouver que 'toute vache est immortelle'.

2.2 Raisonnement par l'absurde

Le principe du raisonnement par l'absurde est le suivant : pour démontrer qu'une proposition R est vraie, on suppose le contraire (c'est-à-dire R fausse), et on essaye d'arriver à un résultat faux (absurde). Par exemple, pour montrer qu'il n'existe pas de plus petit réel strictement positif, on va supposer qu'il en existe un, noté a (donc $0 < a$ est tel qu'il n'existe aucun réel x tel que $0 < x < a$). Or le réel $\frac{a}{2}$ est tel que $0 < \frac{a}{2} < a$, ce qui contredit l'hypothèse. On peut appliquer ce principe par exemple à la proposition $(P \Rightarrow Q)$, notée R . On a vu précédemment que $(P \Rightarrow Q)$ s'écrit $(Q \vee nonP)$ et que $non(Q \vee nonP)$ est équivalente à $((nonQ) \wedge P)$. On peut donc montrer l'implication $(P \Rightarrow Q)$ en montrant que $((nonQ) \wedge P)$ est fausse. Plus précisément on suppose que P est vraie et Q est fausse et l'on démontre que cela aboutit à un résultat faux. Par exemple, pour montrer que, n étant un entier, $(n \text{ impair}) \Rightarrow (\text{le chiffre des unités de } n \text{ est impair})$, on va supposer à la fois que n est impair et que le chiffre de ses unités est pair, ce qui donne : $2m + 1 = 10q + 2p$, donc $1 = 2(5q + p - m)$ ce qui est impossible puisque 1 ne peut pas être le produit de deux entiers dont l'un est 2.

3 Notions sur les ensembles

3.1 Définition d'un ensemble

Un ensemble E est considéré comme une collection d'objets (mathématiques) appelés éléments. $x \in E$ signifie x est un élément de E , $x \notin E$ signifie x n'est pas un élément de E .

Un ensemble E est donc défini si, pour chaque objet x considéré, une et une seule des deux éventualités $x \in E$ et $x \notin E$ est vraie. En pratique, on définit un ensemble, soit en exhibant tous ses éléments, soit en donnant un critère permettant de vérifier la vérité de $(x \in E)$ ou de $(x \notin E)$.

Par exemple l'ensemble des nombres réels positifs ou nuls s'écrit

$$\mathbb{R}^+ = \{x \in \mathbb{R}; x \geq 0\}.$$

Dans la suite, nous supposerons connus les ensembles suivants :

- l'ensemble $\mathbb{N}$ des entiers naturels,
- l'ensemble $\mathbb{Z}$ des entiers relatifs,
- l'ensemble $\mathbb{Q}$ des nombres rationnels,
- l'ensemble $\mathbb{R}$ des nombres réels,
- l'ensemble $\mathbb{C}$ des nombres complexes.
- l'ensemble $\mathbb{R}^*$ des nombres réels non nuls,
- l'ensemble $\mathbb{R}^+$ des nombres réels positifs ou nuls,
- l'ensemble $\mathbb{R}^-$ des nombres réels négatifs ou nuls.

Remarque 3.1 *Le chapitre 2 de ce cours, sera toutefois consacré à un rappel et à certains développements des propriétés fondamentales de $\mathbb{R}$ et de $\mathbb{C}$.*

3.2 Définition d'un sous-ensemble et de l'ensemble vide

Soit E un ensemble, une partie ou sous-ensemble de E est un ensemble A vérifiant la propriété suivante pour tout x : $(x \in A) \Rightarrow (x \in E)$ On dit aussi que A est inclus dans E , et on note $A \subset E$, par exemple $\mathbb{R}^+ \subset \mathbb{R}$. Pour montrer l'égalité de deux ensembles on procède par double inclusion, c'est-à-dire

$$(A = B) \Leftrightarrow (A \subset B) \text{ et } (B \subset A)$$

ou par équivalence, c'est-à-dire

$$(x \in A) \Leftrightarrow (x \in B)$$

qui est la traduction de la double inclusion. L'ensemble vide, noté $\emptyset$ est un ensemble qui ne contient aucun élément, c'est-à-dire qui est tel que la propriété $(x \in \emptyset)$ est fausse quel que soit x . Donc $\emptyset \subset E$ pour tout ensemble E . En effet $(x \in \emptyset)$ étant toujours fausse l'implication $(x \in \emptyset) \Rightarrow (x \in E)$ est vraie.

3.3 Intersection et union d'ensembles

Si A et B sont deux parties de E , on appelle intersection de A et B , notée $A \cap B$ l'ensemble des éléments communs à A et B , et l'on a :

$$(x \in A \cap B) \Rightarrow (x \in A) \text{ et } (x \in B)$$

Par exemple, soit $E = \mathbb{N}$, A l'ensemble des entiers multiples de 3, B l'ensemble des entiers multiples de 5, alors $A \cap B$ est l'ensemble des entiers multiples de 15. De manière générale, si A est l'ensemble des entiers multiples de n et B l'ensemble des entiers multiples de m , alors $A \cap B$ est l'ensemble des entiers multiples du plus petit multiple commun de n et m . (à propos vous souvenez-vous du calcul de ce plus petit multiple commun ?).

On appelle union de A et B , notée $A \cup B$ l'ensemble des éléments appartenant à A ou à B , et l'on a :

$$(x \in A \cup B) \Rightarrow ((x \in A) \text{ ou } (x \in B)).$$

Ainsi, soit $E = \mathbb{N}$, A l'ensemble des entiers multiples de 3, B l'ensemble des entiers multiples de 5, alors $A \cup B$ est l'ensemble des entiers qui sont multiples de 3 ou de 5.

Soit E un ensemble quelconque, pour toutes parties A , B et C de l'ensemble E , on a les égalités ensemblistes suivantes :

$$A \cap (B \cap C) = (A \cap B) \cap C,$$

$$A \cup (B \cup C) = (A \cup B) \cup C$$

$$A \cap (B \cup C) = (A \cap B) \cup (A \cap C),$$

$$A \cup (B \cap C) = (A \cup B) \cap (A \cup C).$$

que l'on peut démontrer par équivalence.

3.4 Complémentaire d'une partie d'un ensemble

Soit E un ensemble, pour toute partie A de E , l'ensemble des éléments de E qui n'appartiennent pas à A s'appelle le complémentaire de A dans E et se note $\mathcal{C}_E A$ ou $E \setminus A$.

Lorsque, du fait du contexte, il n'y a pas d'ambiguïté sur l'ensemble E , on se contente souvent de noter $\mathcal{C}A$, le complémentaire de A dans E .

Par exemple:

- soit $E = \mathbb{N}$ et soit A l'ensemble des entiers pairs, alors $\mathcal{C}A$ est l'ensemble des entiers impairs,

- soit $E = IR$ et soit $A = \{0\}$, alors $\mathcal{C}A = IR^*$
- soit $E = IR$ et soit $B = \{2\}$ alors $\mathcal{C}B = IR \setminus \{2\}$
- soit $E = IR$ et soit $A = IR^+$, alors $\mathcal{C}A = IR_*^-$.

Pour toutes parties A et B d'un ensemble E , on a les propriétés suivantes

$$A \subset B \Rightarrow \mathcal{C}B \subset \mathcal{C}A,$$

$$\mathcal{C}(\mathcal{C}A) = A,$$

$$\mathcal{C}(A \cap B) = \mathcal{C}A \cup \mathcal{C}B,$$

$$\mathcal{C}(A \cup B) = \mathcal{C}A \cap \mathcal{C}B.$$

Notons bien que le complémentaire d'une intersection est l'union des complémentaires et de même le complémentaire d'une union est l'intersection des complémentaires. Notons aussi que lorsque l'on définit un ensemble E comme l'ensemble des éléments vérifiant une propriété P , soit

$$E = \{x, P(x)\}$$

le complémentaire de E est l'ensemble des éléments vérifiant *non* P . De même, si A et B sont définis à l'aide des propriétés P et Q , alors $A \cap B$ est défini par (P et Q) et $A \cup B$ par (P ou Q) .

Attention : la notation $\mathcal{C}_E A$ suppose que A est inclus dans E , on a donc $\mathcal{C}_E(\mathcal{C}_E A) = A$, alors que la notation $B \setminus A$ définie par $x \in B \setminus A \Rightarrow x \in B, x \notin A$ ne suppose pas que A est inclus dans B , on a alors $B \setminus (B \setminus A) = A \cap B$.

Par exemple:

$E = IR, A = [1, 3], B = [2, 5], B \setminus A =]3, 5], B \setminus (B \setminus A) = [2, 3] = A \cap B$, on montrerait de même que $A \setminus (A \setminus B) = [2, 3]$.

3.5 Cardinal d'un ensemble

On dit qu'un ensemble E est fini s'il a un nombre fini d'éléments. Le nombre de ses éléments est appelé cardinal de E et se note $\text{card}(E)$. Par exemple, si $E = \{1, 2, 3\}$, alors $\text{card}(E)=3$. On note $P(E)$ l'ensemble des parties de E .

Proposition 3.1 *Le nombre des parties d'un ensemble fini de cardinal n est égal à 2^n .*

Démonstration- Il suffit de dénombrer le nombre des parties de $E = \{a_1, \dots, a_n\}$:

- partie vide $\emptyset$
- parties ne contenant qu'un élément, il y en a n $\{a_1\}, \dots, \{a_n\}$
- parties formées de deux éléments, il y en a C_2^n de la forme $\{a_i, a_j\}$ avec $i \neq j$,
- . . .
- parties formées de p éléments obtenues en prenant toutes les combinaisons de p éléments parmi n , il y en a C_p^n
- . . .
- E lui-même ($E \subset E$)

Le nombre d'éléments est donc

$$C_0^n + C_1^n + \dots C_p^n + \dots C_n^n = (1 + 1)^n = 2^n$$

Cette dernière relation est une application de la formule du binôme de Newton.

3.6 Produit cartésien d'ensembles

Si E et F sont deux ensembles, le produit cartésien de E par F , noté $E \times F$, est constitué des couples (x, y) , où x décrit E et où y décrit F .

Les couples (x, y) et (x_0, y_0) sont égaux si et seulement si $x = x_0$ et $y = y_0$. Si E est un ensemble, le produit cartésien $E \times E$ se note E^2 . Par exemple, le produit cartésien $\mathbb{R}^2$ est formé des couples de nombres réels ; ceux-ci permettent de déterminer un point du plan par ses coordonnées, lorsqu'on s'est donné un repère (appelé aussi repère cartésien) c'est-à-dire la donnée d'une origine O et de deux vecteurs non colinéaires $(\vec{i}, \vec{j})$.

Plus généralement, pour tout entier n supérieur ou égal à 2, le produit cartésien de E par lui même n fois se note E^n . Les éléments de E^n sont les n -uplets $(x_1, x_2, \dots, x_n)$ où les éléments $x_1, x_2, \dots, x_n$ appartiennent à E . Les n -uplets $(x_1, x_2, \dots, x_n)$ et $(y_1, y_2, \dots, y_n)$ sont égaux si et seulement si on a $x_i = y_i$, pour tout i tel que $1 \leq i \leq n$.

Attention ! lorsque $E \neq F$, $E \times F$ est différent de $F \times E$.

4 Quantificateurs

4.1 Proposition dépendant d'une variable : $P(x)$

Si P est une proposition dont l'énoncé dépend de la valeur d'une variable x on peut la noter $P(x)$ et considérer les cas particuliers $P(a)$ où a est une valeur particulière de x .

Par exemple, soit dans IR la proposition suivante $P(x) : x^2 - 1 < 0$. Alors $P(2)$ est fausse et $P(0)$ est vraie, et plus généralement $P(x)$ est vraie pour toutes les valeurs de x appartenant à $] - 1, +1[$.

Soit E un ensemble, P une propriété dépendant d'une variable x de E , on est souvent amené à considérer l'ensemble des éléments a de E tels que $P(a)$ soit vraie (on dit aussi, les a qui vérifient la propriété P). On le note

$$A_P = \{x \in E, P(x)\}$$

4.2 Quantificateur universel : $\forall$

Pour exprimer qu'une propriété $P(x)$ est vraie pour tous les éléments x d'un ensemble de E , on écrit :

$$\forall x \in E, P(x)$$

(La virgule dans cette phrase' peut être remplacée par un autre signe séparateur, couramment le point-virgule (;) ou le trait vertical (|)). Si on reprend la notation A_P définie précédemment, on a alors :

$$(\forall x \in E, P(x)) \Leftrightarrow A_P = E$$

Exemples

- $\forall x \in IR; x^2 \geq 0$
- $\forall x \in [2, +\infty[, x^2 - 4 \geq 0$
- $E \subset F$ s'écrit : $\forall x \in E; x \in F$

Proposition 4.1 On a l'équivalence suivante :

$$(\forall x \in E, (P(x) \text{ et } Q(x))) \Leftrightarrow ((\forall a \in E, P(a)) \text{ et } (\forall b \in E, Q(b)))$$

Démonstration - En effet $P(x)$ et $Q(x)$ sont vraies pour tout élément de E est bien équivalent à $P(x)$ est vraie pour tout élément de E et $Q(x)$ est vraie pour tout élément de E . On peut également démontrer cette proposition avec les ensembles, en effet :

$$A_P \cap A_Q = E \Leftrightarrow A_P = E \text{ et } A_Q = E$$

4.3 Quantificateur existentiel : $\exists$

Pour exprimer qu'une propriété $P(x)$ est vérifiée pour au moins un élément x d'un ensemble E , on écrit :

$$\exists x \in E, P(x)$$

qui se traduit par : 'il existe x appartenant à E tel que $P(x)$ est vraie'.

S'il existe un unique élément x de E tel que $P(x)$ est vraie, on écrira

$$\exists! x \in E, P(x)$$

Par exemple $\exists x \in IR, x^2 = 1$, mais $\exists! x \in IR; x^2 = 1$ est fausse.

Par contre, vous pouvez écrire

$$\exists! x \in IR^+, x^2 = 1$$

La propriété $(\exists x \in E, P(x))$ se traduit par $A_P \neq \emptyset$.

Proposition 4.2 . On a l'équivalence suivante :

$$\exists x \in E; ((P(x) \text{ ou } Q(x)) \Leftrightarrow (\exists a \in E; P(a)) \text{ ou } (\exists b \in E, Q(b)))$$

Cette proposition peut être démontrée en exercice par double implication.

4.4 Quantificateurs multiples

Il s'agit simplement de successions de $\forall$ et $\exists$. Mais il faut faire très attention à l'ordre dans lequel on les écrit. Par contre on a le résultat évident suivant:

Proposition 4.3 Deux quantificateurs de même nature qui se suivent peuvent être échangés

Par exemple

$$\forall x \in IR, \forall y \in IR; x^2 + y^2 \geq 0 \Leftrightarrow \forall y \in IR, \forall x \in IR; x^2 + y^2 \geq 0$$

De même

$$\exists x \in IR; \exists y \in IR, x + y = 0 \Leftrightarrow \exists y \in IR; \exists x \in IR, x + y = 0$$

Proposition 4.4 On a les équivalences suivantes :

$$\text{non}(\forall x \in E, P(x)) \Leftrightarrow (\exists a \in E, \text{non}P(a))$$

$$\text{non}(\exists a \in E, Q(a)) \Leftrightarrow (\forall x \in E; \text{non}Q(x))$$

Démonstration: La négation de la proposition : (P est vérifiée pour tout élément de E) est: (il existe au moins un élément de E pour lequel la proposition P est fausse). Ce qui s'exprime mathématiquement par :

$$\text{non}(\forall x \in E, P(x)) \Leftrightarrow (\exists a \in E; \text{non}P(a))$$

On peut démontrer cette proposition avec les ensembles :

$$\text{non}(A_P = E) \Leftrightarrow (\exists a \in E, a \notin A_P) \Leftrightarrow (\exists a \in E, \text{non}P(a))$$

La deuxième proposition n'est que la négation de cette première proposition, si l'on appelle Q la proposition ($\text{non}P$).

Par exemple, la négation de $\forall x \in IR, x > 0$ est $\exists a \in IR, a \leq 0$.

Proposition 4.5 Soient E et F deux ensembles et P une proposition dépendant de deux indéterminées et à chaque couple d'éléments (a, b) du produit cartésien $E \times F$, on associe $P(a, b)$.

Alors on a:

$$\text{non}(\forall a \in E, \exists b \in F, P(a, b)) \Leftrightarrow (\exists a \in E, \forall b \in F, \text{non}P(a, b))$$

Démonstration : C'est une application directe de la proposition précédente.

Attention ! L'implication inverse de celle de la proposition ci-dessus est fausse.

Attention ! $\exists x, \forall y$ n'a pas le même sens que $\forall y, \exists x$.

5 Relation binaires

Définition 5.1 . On appelle relation binaire dans E une partie de $E \times E$. Plus précisément, soit $\mathcal{R} \subset E \times E$, on dit que x et y sont liés par la relation $\mathcal{R}$ si $(x, y) \in \mathcal{R}$.

On écrit souvent $x\mathcal{R}y$ pour indiquer que x et y sont liés par la relation $\mathcal{R}$.

Par exemple $E = IN$ et $\mathcal{R} = \{(x, y) \in IN \times IN | x \leq y\}$ définit une partie de $IN \times IN$ et la relation associée est la relation 'inférieur ou égal'. Autre exemple, $E = Z \times Z^*$ et $\mathcal{R} = \{(p, q), (p', q') \in E \times E | pq' = p'q\}$ définit un exemple de relation qui va être appelée relation d'équivalence.

5.1 Relation d'équivalence

Définition 5.2 On appelle relation d'équivalence une relation qui vérifie les trois propriétés suivantes :

- elle est réflexive : $x\mathcal{R}x$
- elle est symétrique : $x\mathcal{R}y \Rightarrow y\mathcal{R}x$
- elle est transitive : $x\mathcal{R}y$ et $y\mathcal{R}z \Rightarrow x\mathcal{R}z$.

Si $\mathcal{R}$ est une relation d'équivalence, on écrit souvent

$x \equiv y$, ou $x \sim y$, au lieu de $x\mathcal{R}y$.

Dans un ensemble quelconque, la relation ' x est égal à y ' est une relation d'équivalence.

À partir d'une relation d'équivalence on peut définir des classes d'équivalence. La relation d'équivalence est d'ailleurs une généralisation de la relation d'égalité. Elle est présente partout en mathématiques. Elle permet, lorsque l'on étudie certains objets mathématiques, de n'en conserver que les propriétés pertinentes pour le problème considéré.

5.2 Classes d'équivalence

Définition 5.3 étant donnée une relation d'équivalence $\mathcal{R}$, on appelle classe d'équivalence de l'élément $a \in E$ la partie $\bar{a} \in E$ définie par :

$$\bar{a} = \{x \in E | x\mathcal{R}a\}$$

Tous les éléments de $\bar{a}$ sont donc équivalents entre eux. On appelle ensemble quotient de E par $\mathcal{R}$, que l'on note $E/\mathcal{R}$, l'ensemble constitué des classes d'équivalences des éléments de E .

L'ensemble des classes d'équivalences constitue une partition de E , c'est-à-dire une famille de sous-ensembles de E deux à deux disjoints et dont la réunion est égale à E .

En effet tout élément $a \in E$ appartient à la classe $\bar{a}$. Par ailleurs s'il existe $x \in \bar{a} \cap \bar{b}$ alors on a $x\mathcal{R}a$ et $x\mathcal{R}b$ et donc, d'après la transitivité, $a\mathcal{R}b$ ce qui implique $\bar{a} = \bar{b}$.

Donc $\bar{a} \neq \bar{b} \Rightarrow \bar{a} \cap \bar{b} = \emptyset$ ce qui correspond bien à une partition.

Soit $E = Z \times Z^*$ et $\mathcal{R} = \{(p, q), (p', q')\} \in E \times E | pq' = p'q\}$, alors la classe d'équivalence de (p, q) avec $q \neq 0$ est l'ensemble des $(p', q'), q' \neq 0$, tels que $\frac{p}{q} = \frac{p'}{q'}$.

On peut alors construire un ensemble, noté Q , comme l'ensemble quotient de $Z \times Z^*$ par la relation d'équivalence précédente. Un élément de Q , qui est appelé nombre rationnel, est donc la classe d'équivalence d'un couple $(p, q), q \neq 0$ et on le note, par abus de langage, $\frac{p}{q}$.

5.3 Relations d'ordre

Définition 5.4 Soit E un ensemble on appelle relation d'ordre , une relation binaire qui vérifie les 3 propriétés suivantes :

- $\mathcal{R}$ est réflexive : $x\mathcal{R}x$
- $\mathcal{R}$ est antisymétrique : si $x\mathcal{R}y$ et $y\mathcal{R}x$ alors $x = y$
- $\mathcal{R}$ est transitive : $x\mathcal{R}y$ et $y\mathcal{R}z \Rightarrow x\mathcal{R}z$.

On dit qu'on a un ordre partiel sur E , et que E est partiellement ordonné par $\mathcal{R}$.

On notera $\leq$ les relations d'ordre partiel, par analogie avec la relation " est plus petit ou égal " qui définit une relation d'ordre sur $\mathbb{N}$.

Définition 5.5 Soit E un ensemble et $\mathcal{R}$ une relation d'ordre sur E . On dit que x et y sont comparables si l'une des affirmations : $x\mathcal{R}y$ ou $y\mathcal{R}x$ est vraie

Définition 5.6 Etant donné un ensemble E et une relation d'ordre $\mathcal{R}$ sur E , on dit que $\mathcal{R}$ est un ordre total sur E si, et seulement si, deux éléments quelconques de E sont comparables.

Un ensemble muni d'une relation d'ordre total est dit totalement ordonné

Définition 5.7 Etant donné un ensemble E et une relation d'ordre $\mathcal{R}$ sur E , on dit que $\mathcal{R}$ est un bon ordre sur E si, et seulement si, toute partie non vide de E possède un plus petit élément. Si $\mathcal{R}$ est un bon ordre sur E , on dit que E est un ensemble bien ordonné par $\mathcal{R}$.

Remarque 5.1 Un ensemble bien ordonné est totalement ordonné. En effet si x et y sont des éléments de E , alors $\{x, y\}$ est une partie de E , elle possède donc un plus petit élément; et donc soit $x \leq y$ soit $y \leq x$.

Mais un ensemble totalement ordonné n'est pas forcément bien ordonné!!

C'est ainsi que $\mathbb{Z}$ et $\mathbb{R}$ sont totalement ordonnés par la relation d'ordre classique $\leq$ sans être bien ordonnés : $\{x \in \mathbb{Z} | x \leq -1\}$ et $]0, 1]$ ne possèdent pas de plus petit élément.

Définition 5.8 Etant donné un ensemble E muni d'une relation d'ordre $\mathcal{R}$, on dit que E est un treillis si, et seulement si, toute partie finie non vide de E admet une borne supérieure et une borne inférieure.

5.4 Eléments particuliers d'un ensemble ordonné

Etant donné un ensemble E muni d'une relation d'ordre $\leq$ et une partie A de E :

- On appelle maximum (ou plus grand élément) de A et on note $Max(A)$ tout élément $x \in A$ vérifiant :

$$(\forall y \in E)(y \in A \Rightarrow y \leq x)$$

S'il existe, le maximum de A est unique.

- On appelle minimum (ou plus petit élément) de A et on note $Min(A)$ tout élément $x \in A$ vérifiant :

$$(\forall y \in E)(y \in A \Rightarrow x \leq y)$$

S'il existe, le minimum de A est unique.

- On appelle majorant de A tout élément $x \in E$ vérifiant:

$$(\forall y \in E)(y \in A \Rightarrow y \leq x)$$

Si un tel élément existe, A est dite majorée par x .

Le maximum est un majorant qui appartient à A .

- On appelle minorant de A tout élément $x \in E$ vérifiant:

$$(\forall y \in E)(y \in A \Rightarrow x \leq y)$$

Si un tel élément existe, A est dite minorée par x .

Le minimum est un minorant qui appartient à A .

- On appelle borne supérieure de A et on note $Sup(A)$ tout élément $x \in E$ vérifiant :

- x est un majorant de A
- si a est un majorant de A : $x \leq a$

Si elle existe, la borne supérieure de A est unique : c'est le minimum de l'ensemble des majorants.

Si la borne supérieure appartient à A , c'est un maximum.

- On appelle borne inférieure de A et on note $Inf(A)$ tout élément $x \in E$ vérifiant :

- x est un minorant de A
- si a est un minorant de A : $a \leq x$

Si elle existe, la borne inférieure de A est unique : c'est le maximum de l'ensemble des minorants.

Si la borne inférieure appartient à A , c'est un minimum.

- On appelle élément maximal de A tout élément $x \in A$ vérifiant :

$$(\forall y \in E)(y \in A \text{ et } x \leq y \Rightarrow x = y)$$

Le maximum est un élément maximal.

- On appelle élément minimal de A tout élément $x \in A$ vérifiant :

$$(\forall y \in E)(y \in A \text{ et } y \leq x \Rightarrow x = y)$$

Le minimum est un élément minimal.

6 Quelques mots sur les applications

Une application d'un ensemble E (dit 'ensemble de départ') dans un ensemble F (dit 'ensemble d'arrivée') est une correspondance qui associe à tout élément de E un et un seul élément de F .

On note cette application :

$$f : E \rightarrow F \text{ ou } f : x \mapsto y = f(x)$$

où : l'élément $y = f(x)$ de F est l'image de x par f et x est l'antécédent de y . Le graphe de f est la partie G de $E \times F$ constituée des éléments de la forme $(x, f(x))$, où $x \in E$.

Une fonction de E dans F est une application de D dans F , où $D \subset E$ est appelé domaine de définition ($D = \text{dom} f$) de la fonction f .

Pour tout ensemble E , l'application de E dans E qui à tout élément x associe x , se note id_E et s'appelle l'application identique ou identité de E .

Deux applications f_1 et f_2 sont égales si elles ont même ensemble de départ E , même ensemble d'arrivée F et si $\forall x \in E, f_1(x) = f_2(x)$.

Définition 6.1 On appelle image d'une application $f : E \rightarrow F$ l'ensemble :

$$\text{Im} f = \{y \in F, \exists x \in E : y = f(x)\}$$

ce qui se traduit par : $\text{Im} f$ est l'ensemble des éléments y de F tels qu'il existe au moins un élément x de E vérifiant $y = f(x)$.

Par exemple, la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$, $f : x \rightarrow x^2$ est définie pour tout réel, c'est donc une application de $\mathbb{R}$ dans $\mathbb{R}$. Montrons par double inclusion que $\text{Im}f = \mathbb{R}^+$.

En effet, l'élément $f(x) = x^2$ est toujours positif d'où $\text{Im}f \subset \mathbb{R}^+$, de plus si $a \in \mathbb{R}^+$ on peut écrire $a = f(\sqrt{a}) (= (\sqrt{a})^2)$, ce qui montre que $\mathbb{R}^+ \subset \text{Im}f$.

Définition 6.2 Soient E et F deux ensembles et $f : E \rightarrow F$ une application. On dit que :

- l'application f est injective si :

$$\forall x \in E, \forall x' \in E (f(x) = f(x')) \Rightarrow (x = x')$$

- l'application f est surjective si :

$$\forall y \in F \exists x \in E \text{ tel que } y = f(x), \text{ (soit } \text{Im}f = F)$$

- l'application f est bijective, si :

$$\forall y \in F; \exists! x \in E \text{ tel que } y = f(x)$$

Par exemple, l'application $x \rightarrow \sin x$ est surjective si on la considère comme application de $\mathbb{R}$ dans $[-1, +1]$. Elle n'est pas injective puisque, par exemple, $\sin 0 = \sin 2\pi$. Elle est bijective si on la considère comme application de $[-\frac{\pi}{2}, \frac{\pi}{2}]$ sur $[-1, +1]$.

Vous montrerez en exercice les résultats suivants :

f n'est pas injective si vous pouvez exhiber deux éléments distincts de E qui ont la même image,

f est injective si et seulement si :

$$\forall x \in E; \forall x' \in E; (x \neq x') \Rightarrow (f(x) \neq f(x'))$$

Proposition 6.1 Soient E et F deux ensembles et $f : E \rightarrow F$ une application. L'application f est bijective si et seulement si f est injective et surjective.

Démonstration - Nous allons raisonner par double implication.

Supposons que l'application est bijective, c'est-à-dire que, quel que soit $y \in F$, il existe un unique $x \in E$ tel que $y = f(x)$. Elle est donc surjective par définition.

Montrons qu'elle est injective.

Soient deux éléments x et x' de E tels que $f(x) = f(x')$. Posons $y = f(x)$ alors on a aussi $y = f(x')$ et comme, par définition de la bijectivité, il n'y a qu'un seul antécédent à y , on a $x = x'$. On a donc $(f(x) = f(x')) \Rightarrow (x = x')$ et f est injective.

Réciproquement, supposons que f est injective et surjective. Puisque f est surjective, $\forall y \in F, \exists x \in E$ tel que $y = f(x)$. Il reste à démontrer que cet antécédent x est unique. Supposons qu'il existe un deuxième antécédent x' vérifiant donc $y = f(x')$, alors l'injectivité $(f(x) = f(x')) \Rightarrow (x = x')$, ce qui montre bien que l'antécédent est unique.

6.1 Composition des applications

Définition 6.3 Soient E, F, G des ensembles et $f : E \rightarrow F, g : F \rightarrow G$ des applications. La composée de f et g , notée $g \circ f$ (ce qui se lit 'g rond f') est l'application de E dans G définie par $(g \circ f)(x) = g(f(x))$ pour tout $x \in E$.

Soit $f : E \rightarrow F$, alors on a $f \circ id_E = f$. En effet, $id_E : E \rightarrow E$ donc la composition $f \circ id_E : E \rightarrow F$ est valide et $\forall x \in E; f \circ id_E(x) = f(id_E(x)) = f(x)$. On montre de la même manière que $id_F \circ f = f$

Proposition 6.2 Soient E, F, G des ensembles et $f : E \rightarrow F, g : F \rightarrow G, h : G \rightarrow H$ des applications, alors on a :

$$h \circ (g \circ f) = (h \circ g) \circ f \text{ (associativité de la composition)}$$

et cette application est notée $h \circ g \circ f : E \rightarrow H$.

Démonstration - Par définition de la composition, on a

$$h \circ (g \circ f) : E \rightarrow H \text{ et } (h \circ g) \circ f : E \rightarrow H,$$

et, pour tout x de E :

$$h \circ (g \circ f)(x) = h((g \circ f)(x)) = h(g(f(x)))$$

$$(h \circ g) \circ f(x) = (h \circ g)(f(x)) = h(g(f(x)))$$

d'où l'égalité des deux applications.

Notons bien que la notation $h \circ g \circ f$ n'a de sens que parce que l'on vient de démontrer que l'on obtient la même application en composant $g \circ f$, puis en formant $h \circ (g \circ f)$ qu'en composant $h \circ g$, pour former enfin $(h \circ g) \circ f$.

Attention à la notation $g \circ f$. Pour former $(g \circ f)(x)$, on applique d'abord f à x , puis g à $f(x)$. La composée est notée $g \circ f$, parce que $(g \circ f)(x) = g(f(x))$.

6.2 Définition de l'application réciproque

Définition 6.4 Soit f une application de E dans F , on dit que g , application de F dans E , est une application réciproque de f si on a

$$f \circ g = id_F \text{ et } g \circ f = id_E \quad (1)$$

C'est à dire

$$\forall y \in F, (f \circ g)(y) = y, \text{ et } \forall x \in E, (g \circ f)(x) = x.$$

Proposition 6.3 Si f admet une application réciproque g , alors cette application réciproque est unique. Lorsque l'application réciproque de f existe on la note f^{-1} .

Démonstration : Pour montrer l'unicité de g , on suppose qu'il existe deux applications g_1 et g_2 qui vérifient (1). Alors on a

$$g_2 = g_2 \circ id_F = g_2 \circ (f \circ g_1) = (g_2 \circ f) \circ g_1 = id_E \circ g_1 = g_1$$

Proposition 6.4 Soit f est une application de E dans F

1. Si f est une application bijective, alors f admet une application réciproque f^{-1} . De plus l'application f^{-1} est bijective de F dans E .
2. Si f admet une application réciproque, alors f est bijective

Démonstration:

1. On suppose que f est bijective de E dans F . La démonstration se déroule en 2 étapes : tout d'abord l'existence de f^{-1} , puis le caractère bijectif de f^{-1} .
 - (a) Comme f est bijective, $\forall y \in F, \exists ! x \in E, y = f(x)$. Puisque x est unique on peut définir une application $g : F \rightarrow E$ par $g(y) = x$. Alors pour tout y de F , on a $f(g(y)) = f(x) = y$, par définition de g , d'où $f \circ g = id_F$. D'autre part, à $x \in E$ on associe $y = f(x)$ et donc $g(y) = x$ (unicité de l'antécédent), ce qui donne $x = g(f(x))$, soit $g \circ f = id_E$. Donc g est l'application réciproque de f .
 - (b) Soit $x \in E$, alors $x = g(f(x))$ et donc il existe $y = f(x)$ élément de F tel que $x = g(y)$, d'où g est surjective. Supposons que $g(y) = g(y')$, alors $f(g(y)) = f(g(y'))$ soit $y = y'$. L'application g est donc injective.

2. On suppose que f admet une application réciproque, on va montrer que f est bijective de E sur F .

$$f(x) = f(x') \Rightarrow f^{-1}(f(x)) = f^{-1}(f(x')) \Rightarrow x = x', \text{ donc } f \text{ est injective.}$$

Soit y quelconque appartenant à F , alors $y = f(f^{-1}(y))$, donc y admet un antécédent dans E : c'est $f^{-1}(y)$. f est surjective sur F

Exemple : la fonction $f : \mathbb{R}^+ \rightarrow \mathbb{R}^+$, $f : x \rightarrow x^2$, est une application bijective (elle ne le serait pas si l'on prenait $\mathbb{R}$ comme ensemble de départ) et l'application inverse est $f^{-1} : \mathbb{R}^+ \rightarrow \mathbb{R}^+$, $f^{-1} : x \rightarrow \sqrt{x}$ n'est autre que la fonction racine carrée !

Vous montrerez en exercice un résultat qui semble évident : $(f^{-1})^{-1} = f$.

6.3 Composition des applications réciproques

Proposition 6.5 Soient E, F, G des ensembles et $f : E \rightarrow F$, $g : F \rightarrow G$ des applications bijectives. Alors l'application $g \circ f : E \rightarrow G$ est bijective et

$$(g \circ f)^{-1} = f^{-1} \circ g^{-1}.$$

Démonstration - Soit $z \in G$, alors $\exists! y \in F$, $z = g(y)$ (g bijective), puis $\exists! x \in E$; $y = f(x)$ (f bijective) et donc $\exists! x \in E$, $z = g(f(x)) = (g \circ f)(x)$, ce qui montre que $g \circ f$ est bijective. Elle admet donc une fonction réciproque $(g \circ f)^{-1} : G \rightarrow E$, unique solution de $(g \circ f)^{-1} \circ (g \circ f) = id_E$ et $(g \circ f) \circ (g \circ f)^{-1} = id_G$. On peut montrer aisément que $f^{-1} \circ g^{-1}$ vérifie ces deux équations, de sorte que l'on a bien : $(g \circ f)^{-1} = f^{-1} \circ g^{-1}$.

Université Mohammed V - Agdal
Faculté des Sciences
Département de Mathématiques et Informatique
Avenue Ibn Batouta, B.P. 1014
Rabat, Maroc

.:: Module Mathématiques I : Algèbre ::.

Filière :

Sciences de Matière Physique (SMP)

et

Sciences de Matière Chimie(SMC)

Chapitre II: Rappels et compléments sur les nombres réels et complexes.

Par

Prof: Jilali Mikram

Groupe d'Analyse Numérique et Optimisation

<http://www.fsr.ac.ma/ANO/>

Email : mikram@fsr.ac.ma

Année : 2005-2006

TABLE DES MATIERES

1	Les nombres réels	3
1.1	Loi de composition interne	3
1.2	Les irrationnels	4
1.3	L'ordre dans $\mathbb{R}$	4
1.4	Partie entière d'un nombre réel $x : E(x)$	5
1.5	Les intervalles ouverts et fermés dans $\mathbb{R}$	5
1.6	Les intervalles dans $\mathbb{R}$	6
1.7	Minorant, majorant	8
1.8	Borne supérieure	8
1.9	Borne inférieure	9
1.10	L'axiome de la borne supérieure	10
2	Les nombres complexes	11
2.1	Lois de composition interne de $\mathbb{R}^2$	11
2.2	Parties réelle et imaginaire d'un nombre complexe	11
2.3	Formule du binôme de Newton	12
2.4	Conjugué et module d'un nombre complexe	13
2.5	Inégalité triangulaire	14
2.6	Argument d'un nombre complexe	15
2.7	Représentation graphique des nombres complexes	16
2.8	La formule de De Moivre	17
2.9	Le théorème de d'Alembert - Gauss	17
2.10	Racines nièmes de l'unité	18
2.11	Racines d'une équation du second degré	19
2.12	Introduction à l'exponentielle complexe	21
2.13	Application au calcul trigonométrique	21

Rappels et compléments sur les nombres réels et complexes

1 Les nombres réels

1.1 Loi de composition interne

Dans l'ensemble des entiers naturels, l'ensemble des entiers relatifs et l'ensemble des rationnels, vous connaissez deux lois de composition interne qui sont les opérations : addition et multiplication. On dit que ce sont des lois de composition interne car si

$$(x, y) \in E \times E (E = \mathbb{N} \text{ ou } \mathbb{Z} \text{ ou } \mathbb{Q})$$

le résultat de l'addition ou de la multiplication est aussi dans E . Ainsi, dans $\mathbb{N}$, on a $\forall n \in \mathbb{N}, \forall m \in \mathbb{N}, n + m \in \mathbb{N}$ et $n * m \in \mathbb{N}$. Par contre la loi "soustraction" n'est pas une loi de composition interne dans $\mathbb{N}$, en effet $2 \in \mathbb{N}, 3 \in \mathbb{N}$ et $2 - 3 \notin \mathbb{N}$. C'est pour remédier à ce défaut que l'on a construit l'ensemble des entiers relatifs $\mathbb{Z}$, où la soustraction est une addition "déguisée" entre le premier nombre et l'opposé du second. L'addition dans $\mathbb{Z}$ vérifie alors les propriétés suivantes :

1. elle est commutative : $\forall a \in \mathbb{Z}, \forall b \in \mathbb{Z}, a + b = b + a$,
2. elle est associative : $\forall a \in \mathbb{Z}, \forall b \in \mathbb{Z}, \forall c \in \mathbb{Z}, a + (b + c) = (a + b) + c$,
3. elle a un élément neutre : $\exists e \in \mathbb{Z}$ tel que $\forall a \in \mathbb{Z}, a + e = e + a = a$
4. tout élément a admet un opposé : $\forall a \in \mathbb{Z}, \exists \hat{a} \in \mathbb{Z}$, tel que $a + \hat{a} = \hat{a} + a = e$

Si une loi de composition interne d'un ensemble E , notée par exemple $\diamond$, vérifie les propriétés (2, 3, 4) ci-dessus, où a, b, c sont maintenant des éléments quelconques de E , et le signe " + " remplacé par " $\diamond$ ", on dit qu'elle est une loi de groupe dans E , ou que $(E, \diamond)$ est un groupe. Si en plus, la loi est commutative, le groupe est dit commutatif ou abélien (du nom du mathématicien norvégien Niels H. Abel (1802-1829)). $(\mathbb{Z}, +)$ est donc un groupe abélien, dont l'élément neutre n'est autre que 0, et l'opposé d'un entier a , noté $-a$, est simplement obtenu en changeant le signe de a . Ces notations sont utilisées couramment pour les groupes dont la loi est notée avec le signe d'addition (" + ") (on les appelle groupes additifs). Dans le cas

contraire, il n'y a pas de raison particulière que l'élément neutre soit 0, et on utilise plutôt le terme "élément inverse" à la place de l'opposé. Q est aussi un groupe commutatif pour la loi addition. Ce n'est pas le cas de $\mathbb{N}$. On remarquera bien la différence essentielle de sens entre $\exists e, \forall a$ (définition de l'élément neutre) et $\forall a, \exists \hat{a}$ (définition de l'opposé).

1.2 Les irrationnels

Les lois "addition" et "multiplication" sont des lois de composition interne dans $\mathbb{R}$ qui possèdent les mêmes propriétés que lorsqu'on les considère dans Q ($Q \subset \mathbb{R}$). Elles donnent donc à $\mathbb{R}$ la structure de corps. Mais l'ensemble $\mathbb{R}$ est plus riche que Q , puisque les rationnels ne permettent pas de représenter tous les nombres "usuels" comme par exemple $\pi, \sqrt{2}, e$ (base du logarithme népérien),...

Un réel qui n'est pas rationnel s'appelle irrationnel. Ainsi, on peut démontrer que $\sqrt{2}$ est un irrationnel.

Raisonnons par l'absurde et supposons que $\sqrt{2} = \frac{p}{q}$ avec $p, q \in \mathbb{N}$, et tels que la fraction est irréductible (c'est-à-dire, les entiers p et q n'ont aucun diviseur commun autre que 1). On a alors $2q^2 = p^2$ donc p^2 est pair donc p est pair, soit $p = 2p'$. On en déduit que $2q^2 = 4p'^2$ donc q^2 est pair donc q est pair, soit $q = 2q'$, ce qui donne $\frac{p}{q} = \frac{2p'}{2q'}$ ce qui est contraire à l'hypothèse d'irréductibilité de la fraction $\frac{p}{q}$.

1.3 L'ordre dans $\mathbb{R}$

Vous savez déjà que deux nombres réels quelconques peuvent être comparés à l'aide de la relation " $\leq$ " définie par :

$$\forall a \in \mathbb{R}, \forall b \in \mathbb{R}, (a \leq b) \Leftrightarrow (a - b \leq 0)$$

Cette relation est une relation d'ordre:

On définit le plus grand (resp. le plus petit) des nombres réels a et b par

$$\max\{a, b\} = \begin{cases} b & \text{si } a \leq b \\ a & \text{si } b \leq a \end{cases} \quad \min\{a, b\} = \begin{cases} b & \text{si } b \leq a \\ a & \text{si } a \leq b \end{cases}$$

Rappelons les propriétés de compatibilité suivantes entre la relation d'ordre " $\leq$ " et les lois d'addition et de multiplication dans $\mathbb{R}$ (qui font de $\mathbb{R}$ un corps ordonné) :

- $\forall a \in \mathbb{R}, \forall b \in \mathbb{R}, \forall c \in \mathbb{R}, (a \leq b) \Rightarrow (a + c \leq b + c),$
- $\forall a \in \mathbb{R}, \forall b \in \mathbb{R}, ((0 \leq a) \text{ et } (0 \leq b)) \Rightarrow (0 \leq ab).$

Notons que la deuxième propriété ci-dessus est équivalente à :

$$\forall a \in \mathbb{R}, \forall b \in \mathbb{R}, \forall c \in \mathbb{R}, (a \leq b) \text{ et } (0 \leq c) \Rightarrow (ac \leq bc)$$

Terminons par la propriété fondamentale suivante dite propriété d'Archimède:
si A est un nombre réel, il existe un entier naturel n tel que $n > A$.

1.4 Partie entière d'un nombre réel $x : E(x)$

Commençons par énoncer le résultat important suivant :

Proposition 1.1 *Soit a un nombre réel strictement positif et soit x un nombre réel, il existe un unique entier $k \in \mathbb{Z}$ tel que*

$$ka \leq x < (k+1)a.$$

En particulier si l'on prend $a = 1$, ceci signifie qu'un nombre réel est toujours compris entre deux nombres entiers relatifs successifs. Par exemple, le réel $3,2$ est tel que $3 \leq 3,2 < 4$ et le réel $-3,2$ est tel que $-4 \leq -3,2 < -3$. On a : $1 \leq \sqrt{2} < 2, \dots$

Si l'on prend maintenant $a = 0,1$ dans la proposition précédente, on a $14a \leq \sqrt{2} < 15a$, $31a \leq \pi < 32a$, $27a \leq e < 28a$.

Définition 1.1 *Soit x un nombre réel, le plus grand entier inférieur ou égal à x s'appelle la partie entière de x , nous le noterons $E(x)$.*

Par exemple, on a $E(3,2) = 3$, $E(-3,2) = -4$, et $E(\sqrt{2}) = 1$.

1.5 Les intervalles ouverts et fermés dans $\mathbb{R}$

Définition 1.2 *Soient a et b deux nombres réels. On appelle intervalle ouvert de $\mathbb{R}$, toute partie de $\mathbb{R}$, ayant l'une des cinq formes ci-dessous :*

1. $\emptyset$ le sous-ensemble vide de $\mathbb{R}$
2. $]a, b[= \{x \in \mathbb{R}, a < x < b\}$
3. $] - \infty, a[= \{x \in \mathbb{R}, x < a\}$
4. $]a, +\infty[= \{x \in \mathbb{R}, a < x\}$
5. $\mathbb{R}$.

Lorsque $a \geq b$, l'intervalle ouvert $]a, b[$, se réduit à la partie vide de $\mathbb{R}$. Lorsque l'on examine l'intersection de deux intervalles ouverts de l'une ou l'autre des formes $]a, b[$, $] - \infty, a[$ ou $]a, +\infty[$, on voit que cette intersection est soit vide, soit à nouveau un intervalle de l'une de ces trois formes, d'où le résultat

Proposition 1.2 *L'intersection d'un nombre fini d'intervalles ouverts est un intervalle ouvert.*

De la même manière nous introduisons la définition suivante

Définition 1.3 *Soient a et b deux nombres réels. On appelle intervalle fermé de $\mathbb{R}$, toute partie de $\mathbb{R}$, ayant l'une des cinq formes ci-dessous :*

1. $\emptyset$ le sous-ensemble vide de $\mathbb{R}$
2. $[a, b] = \{x \in \mathbb{R}, a \leq x \leq b\}$
3. $] - \infty, a] = \{x \in \mathbb{R}, x \leq a\}$
4. $[a, +\infty[= \{x \in \mathbb{R}, a \leq x\}$
5. $\mathbb{R}$.

Enfin, on appelle segment tout intervalle fermé de la forme $[a, b]$ avec $a \leq b$. Remarquons que le segment $[a, a]$ est l'ensemble $\{a\}$ dont le seul élément est a .

Notons tout d'abord qu'un intervalle fermé se réduit à la partie vide lorsque $b < a$. Nous voyons en outre que l'intersection de deux intervalles fermés est soit vide soit un intervalle fermé. Nous avons donc, comme pour les intervalles ouverts, la propriété ci-dessous

Proposition 1.3 *L'intersection d'un nombre fini d'intervalles fermés est un intervalle fermé.*

1.6 Les intervalles dans $\mathbb{R}$

Définition 1.4 *Soient a et b deux nombres réels. On appelle intervalle de $\mathbb{R}$, toute partie de $\mathbb{R}$, ayant l'une des deux formes ci-dessous :*

1. *intervalle ouvert de $\mathbb{R}$,*
2. *intervalle fermé de $\mathbb{R}$,*
3. $]a, b] = \{x \in \mathbb{R}, a < x \leq b\}$,

$$4. [a, b[= \{x \in \mathbb{R}, a \leq x < b\}$$

Lorsque $a < b$, les nombres a et b s'appellent les extrémités des intervalles $[a, b]$, $]a, b[$, $[a, b[$, $]a, b]$, le nombre $b - a$ est la longueur de l'intervalle, le nombre $\frac{a+b}{2}$ est le centre de l'intervalle.

Enfin nous dirons qu'un nombre réel x est compris entre a et b si l'on a $a \leq x \leq b$ dans le cas où $a \leq b$ ou bien si l'on a $b \leq x \leq a$ dans le cas où $b < a$.

Nous allons donner maintenant une propriété caractérisant les intervalles de $\mathbb{R}$.

Proposition 1.4 *Un sous-ensemble J de $\mathbb{R}$ est un intervalle si et seulement si, quels que soient les réels x et y de J , l'intervalle fermé $[x, y]$ est inclus dans J .*

Voici encore deux propositions importantes. La première se démontre très aisément, il suffit de faire une représentation graphique. La deuxième est plus compliquée et repose sur la proposition 1.1

Proposition 1.5 *Soit I un intervalle ouvert, alors quel que soit $x \in I$, il existe un $\alpha > 0$ tel que l'intervalle $]x - \alpha, x + \alpha[$ soit inclus dans I .*

Démonstration : Traitons le cas où I est de la forme $]a, b[$, avec $a < b$. Les autres cas sont plus simples. Nous avons donc : $\forall x \in \mathbb{R}$,

$$(x \in I) \Leftrightarrow (a < x < b)$$

Choisissons un α vérifiant

$$0 < \alpha < \min(x - a, b - x)$$

Il vient alors

$$a = x - (x - a) < x - \alpha < x < x + \alpha < x + (b - x) = b$$

Proposition 1.6 *Dans tout intervalle ouvert non vide, il y a une infinité de nombres rationnels et une infinité de nombres irrationnels.*

Ceci s'écrit souvent "entre deux rationnels il y a au moins un irrationnel et entre deux réels il y a au moins un rationnel".

1.7 Minorant, majorant

Définition 1.5 Soit A une partie non vide de $\mathbb{R}$, on dit que

- A est majorée s'il existe un nombre réel M tel que $\forall x \in A, x \leq M$. Un tel nombre M s'appelle un majorant de A ,
- A est minorée s'il existe un nombre réel m tel que $\forall x \in A, m \leq x$. Un tel nombre m s'appelle un minorant de A ,
- A est bornée si A est majorée et minorée.

Un intervalle $[a, +\infty[$ est minoré, en effet $a, a - 1$ sont des minorants par exemple. Les intervalles $[a, b],]a, b[$ sont bornés, ils sont en effet minorés par a et majorés par b .

Proposition 1.7 Une partie A de $\mathbb{R}$ est bornée si et seulement si il existe un nombre $M \geq 0$ tel que $\forall x \in A, |x| \leq M$.

La démonstration est à faire en exercice.

Remarque 1.1 Un majorant ou minorant de A peut appartenir à A .

Par exemple a est un majorant de $] - 1, a]$, c'est dans ce cas le plus grand élément de l'ensemble. Il existe des cas où c'est impossible. Par exemple $A =] - 1, a[$ n'admet pas de majorant qui appartienne à A .

Démonstration: raisonnons par l'absurde et supposons qu'il existe un majorant M de A ($x \leq M, \forall x \in A$) tel que $M \in A$ ($M < a$). Puisque $M < a$, il existe un réel α tel que $M < \alpha < a$, c'est-à-dire un élément α de A tel que $M < \alpha$, ce qui est absurde puisque M est un majorant de A .

1.8 Borne supérieure

Soit A une partie non vide, majorée de $\mathbb{R}$. Si A possède un plus grand élément, c'est-à-dire s'il existe $a \in A$ tel que $\forall x \in A, x \leq a$, alors a est le plus petit majorant de A . Ceci veut dire que a est un nombre réel ayant les deux propriétés :

- s est un majorant de A ,
- si M est un majorant de A , alors $s \leq M$.

Les deux propriétés de s énoncées ci-dessus n'impliquent pas que s appartienne à A . Il est clair que si un tel plus petit majorant existe, il est unique.

Définition 1.6 *Le plus petit majorant d'une partie A non vide et majorée s'appelle sa borne supérieure et se note $\sup A$.*

Soient a et b deux réels tels que $a < b$, le segment $[a, b]$ et l'intervalle $] - 1, b]$ ont pour plus grand élément b , donc $\sup[a, b] = \sup] - 1, b] = b$.

Lorsque A ne contient pas de plus grand élément (par exemple $A = [a, b[$), l'existence de sa borne supérieure est loin d'être évidente. La proposition suivante donne une caractérisation de la borne supérieure.

Proposition 1.8 *(Caractérisation de la borne supérieure).*

Soit A une partie de $\mathbb{R}$, non vide et majorée, la borne supérieure de A est l'unique réel s tel que

- *si $x \in A$, alors $x \leq s$,*
- *pour tout réel $t < s$, il existe un nombre $x \in A$ tel que $t < x$.*

Démonstration : Soit $t \in \mathbb{R}$ tel que $t < \sup A$, puisque $\sup A$ est le plus petit des majorants de A , on en déduit que t n'est pas un majorant de A . C'est donc qu'il existe un élément $x \in A$ tel que $x > t$. Le nombre réel $s = \sup A$ possède donc les propriétés de la proposition.

Réciproquement, soit s un réel vérifiant les propriétés et démontrons que c'est le plus petit des majorants de A . Tout d'abord la première propriété implique que s est un majorant de A . D'autre part la deuxième montre que si $t < s$, t n'est pas un majorant de A . Tout majorant de A est donc supérieur ou égal à s . Autrement dit s est bien le plus petit des majorants de A .

Exemple : soit $I =] - 1, b[$, alors b est un majorant de I . En utilisant la caractérisation précédente, montrons que $\sup] - 1, b[= b$. En effet, si $t < b$, alors on a $t < \frac{t+b}{2} < b$ et de plus $\frac{t+b}{2} \in I$.

1.9 Borne inférieure

De la même manière que l'on a défini la borne supérieure, on peut donner la définition suivante de la borne inférieure.

Définition 1.7 *Soit A une partie non vide et minorée de $\mathbb{R}$, le plus grand minorant de A s'appelle la borne inférieure et se note $\inf A$.*

On peut alors démontrer la proposition suivante :

Proposition 1.9 *(Caractérisation de la borne inférieure).*

Soit A une partie de $\mathbb{R}$, non vide et minorée, la borne inférieure de A est l'unique réel s tel que

- si $x \in A$, alors $s \leq x$,
- pour tout réel $t > s$, il existe un nombre $x \in A$ tel que $x < t$.

Par exemple, $I = [a, +1[$ admet un plus petit élément a qui est donc la borne inférieure de A , puisque c'est le plus grand des minorants de A (le démontrer en exercice). Et $I =]a, +1[$ admet a comme borne inférieure d'après la proposition précédente.

En effet:

- $\forall x \in I, a < x$,
- pour tout réel $t > a$, il existe un réel $x = \frac{a+t}{2}$ tel que $a < x < t$, et donc tel que $x \in A$ et $x < t$.

1.10 L'axiome de la borne supérieure

Lorsque A ne contient pas de plus grand élément (par exemple $A = [a, b[$), l'existence de sa borne supérieure est loin d'être évidente. En fait, il existe plusieurs constructions équivalentes de $\mathbb{R}$, dont l'une consiste justement à considérer comme un des axiomes de $\mathbb{R}$ la propriété suivante, appelée propriété ou axiome de la borne supérieure:

Axiome de la borne supérieure: Toute partie non vide et majorée de $\mathbb{R}$ admet une borne supérieure.

On peut présenter intuitivement de la façon suivante cette existence de la borne supérieure. Soit s_1 un majorant de A , et a_1 un réel non majorant de A (donc $a_1 < s_1$). Si le milieu de l'intervalle $[a_1, s_1]$ est un majorant de A , appelons le s_2 , et posons $a_2 = a_1$. Si non, on appelle a_2 ce milieu et pose $s_2 = s_1$. On construit ainsi un segment inclus dans $[a_1, s_1]$, de longueur moitié, et dont l'extrémité droite est un majorant de A , et l'extrémité gauche non. Recommencant le processus, qui porte d'ailleurs un nom : la dichotomie, on voit apparaître une succession de segments emboîtés, de longueur chaque fois diminuée de moitié, et dont l'extrémité droite est toujours un majorant de A , et l'extrémité gauche non. Le "point limite" commun des extrémités de ces segments sera la borne cherchée (on verra la définition précise de cette notion de limite plus tard). Mais attention, c'est justement l'existence de ce point limite qui pose problème ! On peut montrer par exemple que dans $\mathbb{Q}$, les segments emboîtés n'ont pas forcément de point limite commun, le même argument ne marche donc pas. La différence entre $\mathbb{Q}$ et $\mathbb{R}$ réside essentiellement dans cette question d'existence de borne supérieure (ou, ce qui revient au même, l'existence de points limites communs des segments emboîtés). On dira pour cela que $\mathbb{R}$ est complet, tandis que $\mathbb{Q}$ ne l'est pas.

L'axiome n'est pas valable dans Q (à faire en exercice).

On a évidemment de même (l'un se déduisant de l'autre) :

Axiome de la borne inférieure : Toute partie non vide et minorée de $\mathbb{R}$ admet une borne inférieure.

2 Les nombres complexes

2.1 Lois de composition interne de $\mathbb{R}^2$

La nécessité d'étendre $\mathbb{R}$ résulte du fait que certaines équations algébriques n'ont pas de racine dans $\mathbb{R}$, la plus célèbre étant $x^2 + 1 = 0$. Mais il y a une différence fondamentale entre le passage de Q à $\mathbb{R}$ et le passage de $\mathbb{R}$ à $\mathbb{C}$. Dans le premier cas, il s'agit d'une extension destinée à remplir l'espace laissé vide entre les rationnels, dans le deuxième cas, il s'agit d'une extension algébrique : on va agrandir l'ensemble en lui rajoutant une composante, la partie imaginaire, pour pouvoir résoudre des équations qui n'ont pas de racines dans $\mathbb{R}$.

Définition 2.1 Sur $E = \mathbb{R}^2$ on définit les deux lois de composition :

- l'addition $(x, y) + (x', y') = (x + x', y + y')$,
- la multiplication $(x, y) \times (x', y') = (xx' - yy', xy' + x'y)$.

Vous montrerez en exercice que l'addition donne à E une structure de groupe commutatif et que la multiplication a les propriétés nécessaires pour que E ait une structure de corps commutatif. Ce corps, noté $\mathbb{C}$, est appelé le corps des nombres complexes. Un nombre complexe, i.e. un élément de $\mathbb{C}$, est donc un couple de réels, obéissant aux lois de composition précédentes.

2.2 Parties réelle et imaginaire d'un nombre complexe

En utilisant les règles de l'addition et de la multiplication ci dessus, on vérifie :

$$(0, 1) \times (0, 1) = (-1, 0)$$

On identifie le nombre complexe $(x, 0)$ (dont la 2^{ème} composante est nulle) au réel x . On note i le nombre complexe $(0, 1)$, on a donc $i^2 = -1$, c'est à dire i est une des racines de l'équation $z^2 + 1 = 0$.

On a d'autre part:

$$(x, y) = (x, 0) + (0, y) = (x, 0) + (0, 1) \times (y, 0)$$

On peut donc écrire un nombre complexe $z = (x, y)$ sous la forme dite canonique : $z = x + iy$. On dit que x est la partie réelle et y la partie imaginaire de z , et on les note respectivement $\Re z$ et $\Im z$:

$$z = x + iy = \Re z + i\Im z$$

Proposition 2.1 *Soient z et z' deux nombres complexes, alors on a*

$$(zz' = 0) \Leftrightarrow ((z = 0) \text{ ou } (z' = 0))$$

Démonstration : - L'implication $\Leftarrow$ est évidente. Réciproquement, supposons que $zz' = 0$.

Alors, soit $z = 0$ et c'est terminé, soit $z \neq 0$ et l'on a

$$z' = \left(\frac{1}{z}\right)z' = \frac{1}{z}(zz') = \frac{1}{z}0 = 0.$$

Cette propriété, qui est triviale dans $\mathbb{R}$ et dans $\mathbb{C}$, n'est pas vraie dans certains ensembles. Par exemple on verra plus tard que l'on peut avoir deux matrices non nulles dont le produit est nul!

2.3 Formule du binôme de Newton

Proposition 2.2 *Pour tous nombres complexes z et z' et pour tout entier $n \geq 2$, on a:*

$$\begin{aligned} (z + z')^n &= z^n + C_1^n z^{n-1} z' + \dots + C_k^n z^{n-k} z'^k + \dots + C_{n-1}^n z z'^{n-1} + z'^n \\ &= \sum_{k=0}^n C_k^n z^{n-k} z'^k \end{aligned}$$

Démonstration: - La formule se démontre par récurrence.

1. Elle est vraie pour $n = 2$ puisque $(z + z')^2 = z^2 + 2zz' + z'^2$ et que $C_1^2 = 2$.
2. Supposons la vraie pour $n - 1$, c'est-à-dire supposons que

$$(z + z')^{n-1} = z^{n-1} + \dots + C_p^{n-1} z^{n-1-p} z'^p + \dots + z'^{n-1}$$

On en déduit que

$$(z + z')^n = (z + z')(z + z')^{n-1} =$$

$$z(z + z')^{n-1} + z'(z + z')^{n-1} = z(z^{n-1} + \dots + C_k^{n-1} z^{n-1-k} z'^k + \dots + z'^{n-1}) +$$

$$+z'(z^{n-1}+\dots+C_{k-1}^{n-1}z^{n-1-(k-1)}z'^{k-1}+\dots+z'^{n-1})=z^n+\dots+(C_k^{n-1}+C_{k-1}^{n-1})z^{n-k}z'^k+\dots+z'^n$$

Calculons, pour $1 \leq k \leq n-1$, la somme :

$$\begin{aligned} C_k^{n-1} + C_{k-1}^{n-1} &= \frac{(n-1)!}{k!(n-1-k)!} + \frac{(n-1)!}{(k-1)!(n-1-k+1)!} \\ &= \frac{(n-1)!}{k!(n-k)!}((n-k)+k) = C_k^n \end{aligned}$$

d'où découle le résultat annoncé.

2.4 Conjugué et module d'un nombre complexe

Définition 2.2 Soit $z = x + iy$ un nombre complexe, alors

- le nombre complexe $x - iy$ s'appelle le conjugué de z et se note $\bar{z}$
- le nombre réel $\sqrt{x^2 + y^2}$ s'appelle le module de z et se note $|z|$

Voici un résumé des principales propriétés des conjugués et des modules :

- $\bar{\bar{z}} = z$, $\overline{(z_1 + z_2)} = \bar{z}_1 + \bar{z}_2$, $\overline{z_1 z_2} = \bar{z}_1 \bar{z}_2$, $\forall z \neq 0$ $\overline{(\frac{1}{z})} = \frac{1}{\bar{z}}$
- $|z|^2 = z\bar{z}$, $|z| = |\bar{z}|$, $|zz'| = |z||z'|$, $|\frac{1}{z}| = \frac{1}{|z|}$
- $\Re z = \frac{1}{2}(z + \bar{z})$, $\Im z = \frac{1}{2i}(z - \bar{z})$, $|z + z'|^2 = |z|^2 + 2\Re(z\bar{z}') + |z'|^2$
- $z = 0 \Leftrightarrow |z| = 0$ et $\forall z \neq 0$, $\frac{1}{z} = \frac{\bar{z}}{|z|^2}$

Démontrons quelques-unes de ces propriétés (vérifier les autres pour être sûr de bien les manipuler):

- Tout d'abord, pour $z = x + iy (\neq 0)$, nous avons $\overline{(\frac{1}{z})} = \frac{1}{\bar{z}}$:

$$\begin{aligned} \frac{1}{z} &= \frac{1}{x + iy} = \frac{x - iy}{(x + iy)(x - iy)} = \frac{x}{x^2 + y^2} - i \frac{y}{x^2 + y^2} \\ \frac{1}{\bar{z}} &= \frac{1}{x - iy} = \frac{x + iy}{(x + iy)(x - iy)} = \frac{x}{x^2 + y^2} + i \frac{y}{x^2 + y^2} \end{aligned}$$

De même, si $z \neq 0$, on a $|\frac{1}{z}| = \frac{1}{|z|}$ puisque

$$|\frac{1}{z}|^2 = (\frac{x}{x^2 + y^2})^2 + (\frac{y}{x^2 + y^2})^2 = \frac{1}{x^2 + y^2} \text{ et } (\frac{1}{|z|})^2 = \frac{1}{x^2 + y^2}$$

Et enfin le calcul de $|z + z'|^2$ s'obtient par

$$|z + z'|^2 = (z + z')(\bar{z} + \bar{z}') = z\bar{z} + z\bar{z}' + z'\bar{z} + z'\bar{z}'$$

Or

$$z\bar{z}' = \overline{z'z}$$

d'où

$$z\bar{z}' + z'\bar{z} = 2\Re(z\bar{z}')$$

de plus

$$z\bar{z} = |z|^2, \quad z'\bar{z}' = |z'|^2$$

de sorte que l'on a bien :

$$|z + z'|^2 = |z|^2 + 2\Re(z\bar{z}') + |z'|^2$$

2.5 Inégalité triangulaire

Proposition 2.3 *Pour tous nombres complexes z et z' , on a :*

- $|\Re z| \leq |z|$ et $|\Im z| \leq |z|$,
- $|z + z'| \leq |z| + |z'|$ (inégalité triangulaire)
- $||z| - |z'|| \leq |z - z'|$.

Démonstration:

- Si $z = x + iy$, alors $|z|^2 = x^2 + y^2$, $|\Re z|^2 = (\Re z)^2 = x^2$ et $|\Im z|^2 = (\Im z)^2 = y^2$, ce qui donne le résultat puisque :

$$\forall a \in \mathbb{R}^+, \forall b \in \mathbb{R}^+, (a^2 \leq b^2) \Leftrightarrow (a \leq b)$$

- De même, l'inégalité triangulaire est équivalente à $|z + z'|^2 \leq (|z| + |z'|)^2$.
Or

$$(|z| + |z'|)^2 - |z + z'|^2 = |z|^2 + 2|z||z'| + |z'|^2 - (|z|^2 + 2\Re(z\bar{z}') + |z'|^2) =$$

$$2(|z||z'| - \Re(z\bar{z}')) = 2(|z\bar{z}'| - \Re(z\bar{z}'))$$

La dernière quantité est positive ou nulle d'après les propriétés des complexes, d'où le résultat.

- La troisième est obtenue en appliquant l'inégalité triangulaire successivement à $z = (z - z') + z'$ et $z' = (z' - z) + z$ Elle vous est laissée à titre d'exercice.

2.6 Argument d'un nombre complexe

Les propriétés des fonctions trigonométriques cosinus et sinus, nous permettent d'affirmer que, étant donnés deux nombres réels a et b vérifiant $a^2 + b^2 = 1$, il existe un angle θ tel que $\cos \theta = a$ et $\sin \theta = b$.

Nous savons aussi que :

$$((\cos \theta = \cos \phi) \text{ et } (\sin \theta = \sin \phi)) \Leftrightarrow (\theta = \phi + 2k\pi, k \in \mathbb{Z})$$

on dit alors que θ est congru à ϕ modulo 2π et on le note $\theta \equiv \phi[2\pi]$. Autrement dit, l'angle θ défini par les équations (II.3.3) n'est défini qu'à $2k\pi$ près. Soit maintenant $z = x + iy$, un nombre complexe non nul, alors on peut l'écrire

$$z = |z| \left(\frac{x}{|z|} + i \frac{y}{|z|} \right)$$

Il existe un θ (défini à $2k\pi$ près) tel que :

$$\cos \theta = \frac{x}{|z|} \text{ et } \sin \theta = \frac{y}{|z|}$$

puisque

$$\left(\frac{x}{|z|} \right)^2 + \left(\frac{y}{|z|} \right)^2 = 1$$

Ceci nous conduit à la définition

Définition 2.3 *Pour tout nombre complexe z différent de 0 le nombre réel θ , défini à $2k\pi$ près, tel que $z = |z|(\cos \theta + i \sin \theta)$ s'appelle l'argument de z et se note $\arg z$.*

Proposition 2.4 *Pour tous nombres complexes z et z' non nuls on a*

$$\arg(zz') \equiv \arg z + \arg z'[2\pi] \text{ et } \arg \frac{1}{z} \equiv -\arg z[2\pi]$$

Démonstration: Soient $z = |z|(\cos \theta + i \sin \theta)$ et $z' = |z'|(\cos \theta' + i \sin \theta')$, alors

$$zz' = |zz'|(\cos \theta \cos \theta' - \sin \theta \sin \theta' + i(\sin \theta \cos \theta' + \cos \theta \sin \theta')) = |zz'|(\cos(\theta + \theta') + i(\sin(\theta + \theta')))$$

d'où la première relation. La deuxième relation est laissée en exercice.

Remarque 2.1 *Il est parfois utile de choisir une détermination particulière de l'argument. Certains auteurs choisissent l'unique θ appartenant à l'intervalle $[0, 2\pi[$, d'autres celui de l'intervalle $]-\pi, +\pi]$. Nous ferons le premier choix et noterons donc $\text{Arg}z (\in [0, 2\pi[)$ cette détermination de l'argument (détermination principale).*

2.7 Représentation graphique des nombres complexes

Nous avons identifié un nombre complexe $z = x + iy$ à un élément (x, y) de $\mathbb{R}^2$, nous pouvons donc représenter ce nombre complexe par un vecteur $\overrightarrow{OM}$ de composantes x et y dans un repère orthonormé $(O, \vec{u}, \vec{v})$. Le nombre z s'appelle l'affixe du point M . Puisque, dans le paragraphe précédent nous avons écrit z sous la forme trigonométrique $z = |z|(\cos \theta + i \sin \theta)$, les composantes du vecteur $\overrightarrow{OM}$ sont donc $|z|\cos\theta$ et $|z|\sin\theta$, ce qui veut dire que $|z|$ représente la longueur du vecteur $\overrightarrow{OM}$ et l'argument θ de z est une mesure de l'angle que fait $\overrightarrow{OM}$ avec le vecteur unitaire $\vec{u}$. Il résulte des opérations que l'on a construites sur $\mathbb{R}^2$ et que l'on a étendues à $\mathbb{C}$ que si z est associé à $\overrightarrow{OM}$, si z' est associé à $\overrightarrow{OM'}$ alors $z + z'$ est associé à $\overrightarrow{OM} + \overrightarrow{OM'}$.

2.8 La formule de De Moivre

Proposition 2.5 *Pour tout nombre réel θ et tout entier $n \in \mathbb{N}$, on a*

$$(\cos \theta + i \sin \theta)^n = \cos n\theta + i \sin n\theta$$

Démonstration: Cette relation se démontre par récurrence:

La formule est évidemment vraie pour $n = 0$ et $n = 1$.

Supposons la vraie pour $n - 1$, c'est-à-dire :

$$(\cos \theta + i \sin \theta)^{n-1} = \cos(n-1)\theta + i \sin(n-1)\theta$$

, et démontrons la pour n . Il vient :

$$\begin{aligned} (\cos \theta + i \sin \theta)^n &= (\cos \theta + i \sin \theta)^{n-1} (\cos \theta + i \sin \theta) = (\cos(n-1)\theta + i \sin(n-1)\theta) (\cos \theta + i \sin \theta) = \\ &= (\cos(n-1)\theta \cos \theta - \sin(n-1)\theta \sin \theta) + i(\cos(n-1)\theta \sin \theta + \sin(n-1)\theta \cos \theta) = \cos n\theta + i \sin n\theta \end{aligned}$$

2.9 Le théorème de d'Alembert - Gauss

Le théorème suivant, de d'Alembert - Gauss, montre que $\mathbb{C}$ permet de résoudre certaines équations algébriques :

théorème 2.1 *Toute équation algébrique dans $\mathbb{C}$, c'est-à-dire toute équation de la forme*

$$a_n z^n + a_{n-1} z^{n-1} + \dots + a_0 = 0 \quad (II.3.4)$$

où les coefficients $a_i, 0 \leq i \leq n$ sont des nombres complexes, $n \geq 1$ et $a_n \neq 0$, admet au moins une racine z dans $\mathbb{C}$.

Corollaire 2.1 *L'équation (II.3.4) admet exactement n racines dans $\mathbb{C}$ (en comptant chaque racine multiple autant de fois que sa multiplicité).*

La démonstration du théorème sort du cadre de ce cours, par contre on verra (au chapitre sur les polynômes) que le corollaire est tout à fait accessible (si l'on admet le théorème, bien entendu).

Par exemple, l'équation $z^2 + 1 = 0$ admet pour racines les nombres complexes $z_1 = i$ et $z_2 = -i$.

Les paragraphes suivants permettent d'obtenir les racines dans certains cas particuliers.

2.10 Racines nièmes de l'unité

Étant donné un nombre complexe α non nul, on va chercher tous les nombres complexes z possibles vérifiant $z^n = \alpha$. Ces nombres complexes seront appelés les racines nièmes de α . On démontre que tout nombre complexe non nul admet exactement n racines nièmes.

Proposition 2.6 Soit $n \in \mathbb{N}$ tel que $n \geq 2$ et $\alpha \in \mathbb{C}$ non nul. Alors

$$(z^n = \alpha) \Leftrightarrow |z| = \sqrt[n]{|\alpha|} \text{ et } \text{Arg}z = \frac{\text{Arg}\alpha}{n} + \frac{2k\pi}{n}, \text{ où } 0 \leq k \leq n-1$$

Démonstration: a/ ($\Rightarrow$) Si $z^n = \alpha$, alors $|z|^n = |z^n| = |\alpha|$, d'où $|z| = \sqrt[n]{|\alpha|}$ et aussi $\text{Arg}z^n = \text{Arg}\alpha$, ce qui donne (proposition II.3.4) $n\text{Arg}z \equiv \text{Arg}\alpha[2\pi]$. Il existe donc $k \in \mathbb{Z}$ tel que $n\text{Arg}z = \text{Arg}\alpha + 2k\pi$. Les inégalités

$$0 \leq \text{Arg}z < 2\pi \text{ et } 0 \leq \text{Arg}\alpha < 2\pi$$

donnent aisément $0 \leq k \leq n-1$, d'où le résultat.

b/ ($\Leftarrow$) Supposons les relations de droite vérifiées, alors

$$\cos(n\text{Arg}z) = \cos \text{Arg}\alpha \text{ et } \sin(n\text{Arg}z) = \sin \text{Arg}\alpha$$

et d'après la formule de De Moivre, il vient

$$\begin{aligned} z^n &= |z|^n (\cos \text{Arg}z + i \sin \text{Arg}z)^n \\ &= |\alpha| (\cos n\text{Arg}z + i \sin n\text{Arg}z) \\ &= |\alpha| (\cos \text{Arg}\alpha + i \sin \text{Arg}\alpha) \\ &= \alpha \end{aligned}$$

Un cas particulier important est celui des racines nièmes de l'unité. Elles sont solution de $z^n = 1$ et correspondent à $\alpha = 1$. On obtient donc

$$|z| = 1 \text{ et } \text{Arg}z = \frac{0}{n} + \frac{2k\pi}{n}, 0 \leq k \leq n-1$$

soit les racines suivantes :

$$z_k = \cos \frac{2k\pi}{n} + i \sin \frac{2k\pi}{n}, k = 0, 1, \dots, n-1$$

Les racines de l'unité étant de module 1, elles sont représentées graphiquement sur le cercle de rayon 1 et de centre O .

Remarque importante: La définition des racines d'un nombre complexe est une extension stricte du cas réel. Si $a \in \mathbb{R}$ est strictement positif,

on appelle habituellement racine carrée de a le nombre positif r tel que $r^2 = a$. En fait, si l'on note par $\sqrt{a}$ ce nombre r , le nombre $r' = -\sqrt{a}$ a aussi son carré égal à a , donc est une racine de a au sens de la définition ci-dessus. C'est par convention, que l'on dit que "dans $\mathbb{R}$, le nombre positif $\sqrt{a}$ est la racine de a ", même si, "dans $\mathbb{C}$, a admet deux racines, les nombres $\sqrt{a}$ et $(-\sqrt{a})$ ", toutes deux réelles !

Si $a \in \mathbb{C}$ (non réel), alors $\sqrt{a}$ n'a pas de sens puisque le nombre complexe a a deux racines carrées et qu'il n'existe pas dans ce cas de convention pour privilégier l'une ou l'autre. Ceci est précisé dans la proposition suivante:

Proposition 2.7 *Tout nombre complexe non réel z admet exactement deux racines carrées, qui sont opposées.*

Démonstration: Soit $z = a + ib$; $b \neq 0$ et $r = x + iy$;

$$z = r^2 \Leftrightarrow \begin{cases} x^2 - y^2 &= a \\ 2xy &= b \end{cases} \Leftrightarrow \begin{cases} x^2 - y^2 &= a \\ 2xy &= b \\ x^2 + y^2 &= |z| \end{cases}$$

$$\Leftrightarrow \begin{cases} x^2 &= \frac{|z|+a}{2} \\ y^2 &= \frac{|z|-a}{2} \\ 2xy &= b \end{cases} \Leftrightarrow \begin{cases} x &= \varepsilon \sqrt{\frac{|z|+a}{2}} \\ y &= \varepsilon' \sqrt{\frac{|z|-a}{2}} \\ \varepsilon, \varepsilon' &\in \{-1, 1\}, \varepsilon \varepsilon' \text{ du signe de } b \end{cases}$$

2.11 Racines d'une équation du second degré

Soient a, b, c trois nombres complexes, on suppose $a \neq 0$, on recherche les nombres complexes z qui vérifient

$$az^2 + bz + c = 0$$

Ceci va généraliser ce que l'on sait faire lorsque les coefficients a, b, c sont réels. On peut d'ailleurs faire un raisonnement semblable.

$$\begin{aligned} az^2 + bz + c &= a\left(z + \frac{b}{2a}\right)^2 + c - \frac{b^2}{4a} \\ &= a\left(\left(z + \frac{b}{2a}\right)^2 - \frac{b^2 - 4ac}{4a^2}\right) \end{aligned}$$

On définit le nombre complexe $\Delta = b^2 - 4ac$.

Si $\Delta = 0$, alors $az^2 + bz + c = a\left(z + \frac{b}{2a}\right)^2$ ce qui implique que $-\frac{b}{2a}$ est racine double de l'équation.

Si $\Delta \neq 0$, si on note r_0 et r_1 les deux racines carrées (complexes) de Δ , alors $\frac{r_0}{2a}$, $\frac{r_1}{2a}$ sont les deux racines carrées de $\frac{b^2-4ac}{4a^2}$, on a donc :

$$\begin{aligned} az^2 + bz + c = 0 &\Leftrightarrow \left(z + \frac{b}{2a}\right)^2 = \frac{\Delta}{4a^2} \\ &\Leftrightarrow z + \frac{b}{2a} = \frac{r_0}{2a} \text{ ou } z + \frac{b}{2a} = \frac{r_1}{2a} \\ &\Leftrightarrow z = \frac{-b + r_0}{2a} \text{ ou } z = \frac{-b + r_1}{2a} \end{aligned}$$

Montrer en exercice que dans le cas a, b, c réels, on retrouve les formules que vous connaissez.

Exemple: Résoudre l'équation du second degré :

$$z^2 - iz + 1 - 3i = 0$$

on a

$$\Delta = -5 + 12i$$

on cherche $r = x + iy$ tel que $r^2 = \Delta = -5 + 12i$

D'après la proposition (2.7) on a :

$$\begin{aligned} &\begin{cases} x &= \varepsilon \sqrt{\frac{|\Delta|-5}{2}} \\ y &= \varepsilon' \sqrt{\frac{|\Delta|+5}{2}} \\ \varepsilon, \varepsilon' &\in \{-1, 1\}, \quad \varepsilon\varepsilon' > 0 \end{cases} \\ &\Leftrightarrow \begin{cases} x &= 2 \text{ et } y = 3 \\ &\text{ou} \\ x &= -2 \text{ et } y = -3 \end{cases} \quad \text{car } |\Delta| = \sqrt{25 + 144} = 13 \end{aligned}$$

on obtient $r_1 = 2 + 3i$ et $r_2 = -2 - 3i$, et donc

$$z_1 = \frac{i + r_1}{2} = \frac{4i + 2}{2} = 1 + 2i$$

et

$$z_2 = \frac{i + r_2}{2} = \frac{-2i - 2}{2} = -1 - i$$

2.12 Introduction à l'exponentielle complexe

Il est commode de poser

$$e^{i\theta} = \cos \theta + i \sin \theta$$

Cette notation dite "exponentielle complexe", a priori curieuse, est justifiée par le fait qu'elle entraîne les règles opératoires qui rappellent les fonctions de l'exponentielle réelle.

En effet, vous montrerez en exercice que

$$e^{i\theta_1} e^{i\theta_2} = e^{i(\theta_1 + \theta_2)}, e^{i0} = 1, \frac{1}{e^{i\theta}} = e^{-i\theta}$$

Remarquons que

$$\overline{e^{i\theta}} = \overline{\cos \theta + i \sin \theta} = \cos \theta - i \sin \theta = \cos(-\theta) + i \sin(-\theta) = e^{-i\theta}$$

Cette notation permet d'écrire un nombre complexe donné par son module ρ et son argument θ sous la forme simplifiée

$$z = \rho e^{i\theta}$$

Ainsi la formule de De Moivre s'écrit

$$z^n = (\rho e^{i\theta})^n = \rho^n e^{in\theta}, n \in \mathbb{N}$$

Les formules d'Euler expriment $\cos \theta$ et $\sin \theta$ à l'aide de l'exponentielle complexe :

$$\cos \theta = \frac{e^{i\theta} + e^{-i\theta}}{2}, \sin \theta = \frac{e^{i\theta} - e^{-i\theta}}{2i}$$

Attention ! $e^{i\theta_1} = e^{i\theta_2}$ n'implique pas que $\theta_1 = \theta_2$ mais que $\theta_1 = \theta_2 + 2k\pi, k \in \mathbb{Z}$.

2.13 Application au calcul trigonométrique

L'utilisation directe de la formule de De Moivre permet d'exprimer $\cos n\theta$ et $\sin n\theta$ en fonction des puissances de $\cos \theta$ et $\sin \theta$, lorsque l'on utilise la formule du binôme de Newton.

Par exemple, on a $(\cos \theta + i \sin \theta)^3 = \cos 3\theta + i \sin 3\theta$, et la formule du binôme de Newton donne

$$(\cos \theta + i \sin \theta)^3 = \cos^3 \theta + 3i \cos^2 \theta \sin \theta - 3 \cos \theta \sin^2 \theta - i \sin^3 \theta$$

d'où

$$\cos 3\theta = \cos^3\theta - 3\cos\theta\sin^2\theta$$

$$\sin 3\theta = 3\cos^2\theta\sin\theta - \sin^3\theta$$

Mais ce qui est le plus utile c'est l'inverse ! c-à-d de pouvoir exprimer les puissances de $\cos\theta$ et $\sin\theta$ en expression linéaire de $\cos k\theta$ et $\sin k\theta$, par exemple pour pouvoir les intégrer (voir le cours sur les intégrales). On peut alors utiliser l'exponentielle complexe. Ainsi

$$\cos n\theta = \frac{1}{2^n}(e^{i\theta} + e^{-i\theta})^n$$

$$\sin n\theta = \frac{1}{2^n}(e^{i\theta} - e^{-i\theta})^n$$

On développe alors par le binôme de Newton et on regroupe les termes $e^{ik\theta}$ et $e^{-ik\theta}$. Illustrons par un exemple. Choisissons $n = 4$ et appliquons la méthode précédente :

$$\begin{aligned}\cos^4\theta &= \frac{1}{2^4}(e^{i\theta} + e^{-i\theta})^4 \\ &= \frac{1}{16}(e^{i4\theta} + 4e^{i2\theta} + 6 + 4e^{-i2\theta} + e^{-i4\theta}) \\ &= \frac{1}{16}(2\cos 4\theta + 8\cos 2\theta + 6) \\ &= \frac{1}{8}\cos 4\theta + \frac{1}{2}\cos 2\theta + \frac{3}{8}\end{aligned}$$

Université Mohammed V - Agdal
Faculté des Sciences
Département de Mathématiques et Informatique
Avenue Ibn Batouta, B.P. 1014
Rabat, Maroc

.:: Module Mathématiques I : Algèbre ::.

Filière :

Sciences de Matière Physique (SMP)
et
Sciences de Matière Chimie(SMC)

Chapitre III: Polynômes sur $\mathbb{R}$ et $\mathbb{C}$

Chapitre IV: Fractions rationnelles

Par

Prof: Jilali Mikram
Groupe d'Analyse Numérique et Optimisation
<http://www.fsr.ac.ma/ANO/>
Email : mikram@fsr.ac.ma

Année : 2005-2006

TABLE DES MATIERES

1	Polynômes	3
1.1	Présentation des polynômes	3
1.2	Lois sur $\mathbb{K}[X]$	4
1.3	Division euclidienne	5
1.4	Zéros d'un polynôme	7
1.5	Polynôme dérivé	7
1.6	Formules de Mac-Laurin et de Taylor	10
1.7	Ordre de multiplicité d'une racine	11
1.8	Théorème de d'Alembert	12
1.9	Division suivant les puissances croissantes	13
2	Fractions rationnelles	14
2.1	Degré, partie entière	14
2.2	Pôles et partie polaires	15
2.3	Décomposition en éléments simples	17
2.4	Pratique de la décomposition en éléments simples	19

Chapitre 1

Polynômes

1.1 Présentation des polynômes

Définition 1.1.1 On se place sur un corps commutatif $\mathbb{k}$. Un polynôme est défini par la donnée de ses coefficients $a_0, \dots, a_n$ éléments de $\mathbb{k}$. X étant une lettre muette, on note $P(X) = a_0 + a_1X + \dots + a_nX^n$ ou $\sum_{k \geq 0} a_k X^k$, étant entendu que la somme ne comporte qu'un nombre fini de a_k non nuls.

On distingue parfois le polynôme $P(X)$ (qui, par construction, est nul si et seulement si tous ses coefficients sont nuls (*)) de la fonction polynômiale associée:

$$\begin{aligned} P : \mathbb{k} &\longrightarrow \mathbb{k} \\ x &\longrightarrow a_0 + a_1x + \dots + a_nx^n = P(x) \end{aligned}$$

Celle-ci est nulle si et seulement si : $\forall x \in \mathbb{k}, P(x) = 0$ (**)

On a bien évidemment l'implication :

$$P(X) = 0 \implies \forall x \in \mathbb{k}, P(x) = 0.$$

Mais la réciproque est loin d'être évidente. Nous allons montrer que, lorsque $\mathbb{k}$ est égal à $\mathbb{R}$ ou $\mathbb{C}$, il y a équivalence, ce qui permet de confondre polynôme et fonction polynômiale. La phrase $P = 0$ gardera cependant de préférence le sens (*).

Proposition 1.1.1 :

i) Soit P un polynôme à coefficients dans $\mathbb{R}$ ou $\mathbb{C}$. Alors si la fonction polynômiale associée à P est identiquement nulle. P a tous ses coefficients nuls.

ii) Soit P et Q deux polynômes dans $\mathbb{R}$ ou $\mathbb{C}$. Alors, si les fonctions polynomiales associées sont égales (prennent les mêmes valeurs), les deux polynômes sont égaux (ont leurs coefficients égaux).

Démonstration:

i) $\mathbb{k}$ contenant $\mathbb{R}$, nous supposons que la variable x ne prend que des valeurs dans $\mathbb{R}$. Soit

$$P = \sum_{k \geq 0} a_k X^k \text{ tel que } \forall x \in \mathbb{k}, P(x) = 0$$

Alors, pour $x = 0$, on obtient $a_0 = 0$. Donc:

$$\forall x \in \mathbb{R}, a_1 x + \dots + a_n x^n = 0$$

$$\implies \forall x \neq 0, a_1 + \dots + a_n x^{n-1} = 0$$

On ne peut plus prendre $x = 0$, cependant, on peut prendre la limite lorsque x tend vers 0, ce qui donne $a_1 = 0$, etc...

ii) se prouve en appliquant i) à $P - Q$.

Si $P \neq 0$, on appelle degré de P le maximum des k , tels que $a_k \neq 0$. Si $P = 0$, on pose $\deg(P) = -\infty$

Si P est de degré n , $a_n X^n$ est le terme (ou monôme) dominant. Si $a_n = 1$, le polynôme est dit unitaire ou normalisé.

On note $\mathbb{k}[X]$ l'ensemble des polynômes sur le corps $\mathbb{k}$.

1.2 Lois sur $\mathbb{k}[X]$

a) Somme de deux polynômes:

$$\text{Si } P = \sum_{k \geq 0} a_k X^k \text{ et } Q = \sum_{k \geq 0} b_k X^k, \text{ alors } P + Q = \sum_{k \geq 0} (a_k + b_k) X^k$$

On a :

$\deg(P + Q) \leq \max(\deg(P), \deg(Q))$. L'égalité a lieu si les polynômes sont de degrés différents, ou s'ils sont de même degré et que les termes de plus haut degré ne s'éliminent pas.

b) Produit de deux polynômes:

$$\text{Si } P = \sum_{k \geq 0} a_k X^k \text{ et } Q = \sum_{k \geq 0} b_k X^k, \text{ alors } PQ = \sum_{k \geq 0} \sum_{i=0} a_i b_{k-i} + b_k X^k$$

On a :

$$\deg(P.Q) = \deg(P) + \deg(Q)$$

Remarque

Si $PQ = 0$ alors $P = 0$ ou $Q = 0$.

c) Produit d'un polynôme par un scalaire:

$$\text{Si } P = \sum_{k \geq 0} a_k X^k, \text{ alors } \lambda P = \sum_{k \geq 0} \lambda a_k X^k$$

On note $\mathbb{K}_n[X] = \{P \in \mathbb{K}[X] \mid \deg(P) \leq n\}$.

1.3 Division euclidienne

(Ou division suivant les puissances décroissantes)

Donnons un exemple :

$$\begin{array}{r|l} 2X^4 + X^3 - X^2 + X + 1 & 2X^2 - X - 2 \\ 2X^4 - X^3 - 2X^3 & \hline X^2 + X + 1 & \\ \hline 2X^3 + X^2 + X + 1 & \\ 2X^3 - X^2 - 2X & \\ \hline 2X^2 + 3X + 1 & \\ 2X^2 - X - 2 & \\ \hline 4X + 3 & \end{array}$$

Nous affirmons alors que :

$$\underbrace{2X^4 + X^3 - X^2 + X + 1}_{\text{Dividende}} = \underbrace{(2X^2 - X - 2)}_{\text{diviseur}} \underbrace{(X^2 + X + 1)}_{\text{quotient}} + \underbrace{(4X + 3)}_{\text{Reste}}$$

Ce résultat est général :

Proposition 1.3.1 Soit A et B deux polynômes tel que $B \neq 0$. Alors il existe un unique couple (Q, R) tels que :

$$A = BQ + R, \text{ avec } \deg(R) < \deg(B)$$

Q est le quotient, R est le reste.

Lorsque le reste est nul, on dit que B divise A . Un polynôme qui n'est divisible que par lui même (à une constante multipliative près) ou par les constantes est dit irréductible. Par exemple , $X - 3$ dans $\mathbb{C}$, ou $X^2 + 1$ dans $\mathbb{R}$.

On notera l'analogie dans l'énoncé avec la division euclidienne dans $\mathbb{Z}$. Les démonstrations; en ce qui concerne l'unicité, sont également analogues.

Démonstration:

Montrons l'unicité :

Si $A = BQ + R = BQ' + R'$ avec $\deg(R) < \deg(B)$ et $\deg(R') < \deg(B)$, on a $B(Q - Q') = R' - R$; avec $\deg(B(Q - Q')) = \deg(B) + \deg(Q - Q')$ et $\deg(R - R') \leq \max(\deg(R), \deg(R')) < \deg(B)$.

Il ne peut y avoir égalité que si $Q - Q' = 0$ et alors $R - R' = 0$.

Montrons l'existence:

Supposons que $\deg(A) = n$ et $\deg(B) = p$

1^{er} cas: si $n < p$ alors on prend $Q = 0$ et $R = A$.

2^{ème} cas: si $n \geq p$.

On procède par récurrence sur le degré n de A .

Supposons que la propriété est vraie jusqu'à $n - 1$.

Soit $A = a_n x^n + \dots + a_0$ et $B = b_p x^p + \dots + b_0$

On définit $A' = A - \frac{a_n}{b_p} x^{n-p} B$.

A' est degré $n - 1$ et donc d'après l'hypothèse de récurrence on a :

$$A' = BQ' + R' \text{ avec } \deg(R') < \deg(B)$$

$$\text{Or } A = A' + \frac{a_n}{b_p} x^{n-p} B$$

$$= BQ' + R' + \frac{a_n}{b_p} x^{n-p} B$$

$$= B(Q' + \frac{a_n}{b_p} x^{n-p}) + R'$$

$$= BQ + R' \text{ avec } \deg R' < \deg B \quad \text{C.Q.F.D}$$

1.4 Zéros d'un polynôme

Définition 1.4.1 On dit que a , élément de $\mathbb{K}$, est un zéro ou une racine du polynôme P si a annule la fonction polynomiale associée à P , c'est à dire $P(a) = 0$.

On a alors le résultat suivant:

Proposition 1.4.1 :

a est un zéro de P si et seulement si P est divisible par $X - a$.

Démonstration:

Si P est divisible par $X - a$, alors il existe Q tel que $P(X) = (X - a)Q(X)$.

On a alors $P(a) = 0$.

Réciproquement, si $P(a) = 0$, considérons la division euclidienne de P par $X - a$. On a :

$P(X) = (X - a)Q(X) + R$ avec $\deg(R) < \deg(X - a) = 1$, donc R est une constante. On obtient alors

$0 = P(a) = R$ donc $R = 0$ et P est divisible par $X - a$.

Une autre démonstration consiste à écrire que, si $P(X) = \sum_{k \geq 0} a_k X^k$ et si $P(a) = 0$ alors

$$P(X) - P(a) = \sum_{k \geq 0} a_k (X^k - a^k)$$

dont chaque terme se factorise par $X - a$.

Il se peut que P se factorise par une puissance de $X - a$. Si k est la puissance maximale de $X - a$ par laquelle le polynôme P se factorise (de sorte que $P = (X - a)^k Q$ avec $Q(a) \neq 0$), on dit que k est l'ordre de multiplicité de la racine a .

1.5 Polynôme dérivé

On définit le polynôme dérivé de $P = \sum_{k \geq 0} a_k X^k$ comme étant égal à

$$P' = \sum_{k \geq 1} k a_k X^{k-1}.$$

On peut définir de la même façon les dérivées successives.

P.G.C.D, et ALGORITHME d'EUCLIDE.

Définition 1.5.1 *On dit qu'un polynôme B divise un polynôme A s'il existe un polynôme Q tel que $A = BQ$.*

On dit alors que B est un diviseur de A ou encore que A est un multiple de B .

Diviseurs communs à deux polynômes.

On a souvent à résoudre le problème suivant: étant donné deux polynômes A et B , quels sont leurs diviseurs communs? Ceci peut servir, par exemple, à trouver les racines communs aux deux équations $A(x) = 0$ et $B(x) = 0$, ou à simplifier la fraction rationnelle A/B .

Pour résoudre ce problème on utilise à répétition le lemme suivant:

Lemme: Les diviseurs communs à deux polynômes A et B sont les mêmes que les diviseurs communs à B et R où R est le reste de la division de A par B .

Démonstration :

Soit $A = BQ + R$ avec $\deg R < \deg B$.

Soit C un diviseur commun à A et B .

on a $A = CQ_1$ et $B = CQ_2$.

$$\begin{aligned} R &= A - BQ \\ &= CQ_1 - CQ_2Q \\ &= C(Q_1 - Q_2Q) \end{aligned}$$

Donc C est un diviseur commun à B et R

Réciproquement :

Soit C un diviseur commun à B et R

On a : $B = CQ_1$ et $R = CQ_2$

Comme $A = BQ + R$

$$\begin{aligned} \text{Donc } A &= CQ_1Q + CQ_2 \\ &= C(Q_1Q + Q_2) \end{aligned}$$

Donc C est un diviseur commun à A et B .

Algorithme d'Euclide:

C'est le procédé qui consiste à répéter le lemme précédent. Pour l'exposer, il est commode de changer un peu les notations. Mais faisons tout d'abord la remarque suivante:

Remarque: Les diviseurs communs à A et 0 sont les diviseurs de A .

Cherchons maintenant les diviseurs communs à A_0 et A_1 , non nuls, avec $\deg(A_0) \geq \deg(A_1)$. Par divisions successives, on peut écrire:

$$A_0 = Q_1 A_1 + A_2 \text{ avec } \deg(A_2) < \deg(A_1) \text{ ou } A_2 = 0.$$

Puis si $A_2 \neq 0$:

$$A_1 = Q_2 A_2 + A_3 \text{ avec } \deg(A_3) < \deg(A_2) \text{ ou } A_3 = 0.$$

Si l'on continue ainsi, on définit des polynômes A_0, A_1 etc... avec $\deg(A_0) \geq \deg(A_1) \geq \deg(A_2) \geq \deg(A_3) \dots$ et ceci ne peut se poursuivre indéfiniment, ce qui implique que l'on arrive à un reste $A_k = 0$. Les dernières divisions sont donc:

$$A_{k-3} = Q_{k-2} A_{k-2} + A_{k-1} \text{ avec } \deg(A_{k-1}) < \deg(A_{k-2})$$

$$A_{k-2} = Q_{k-1} A_{k-1} \text{ (c'est à dire } A_k = 0).$$

Les diviseurs communs à A_0 et A_1 sont les diviseurs de A_{k-1} , ainsi on a le théorème suivant:

Théorème 1.5.1 *étant donnés deux polynômes A et B , il existe un polynôme D tel que les diviseurs communs à A et B soient les diviseurs de D . Ce polynôme est le dernier reste non nul dans la suite des diviseurs de l'algorithme d'Euclide.*

Définition 1.5.2 *Le polynôme D du théorème précédent est appelé PGCD de A et B , et on note $\text{PGCD}(A, B) = D$ ou $A \wedge B = D$.*

Remarques:

a) Lorsque D est une constante, les polynômes A et B n'ont pas d'autres diviseurs communs que les constantes non nulles.

b) le PGCD n'est défini qu'à une constante multiplicative près (on choisit donc le polynôme unitaire correspondant).

Si $D = 1$, on dit que A et B sont premiers entre eux.

On peut également définir le PGCD de n polynômes. Si celui-ci vaut 1, ces polynômes sont dits premiers entre eux dans leur ensemble (exemple: $(X+1)(X+2), (X+1)(X+3), (X+2)(X+3)$).

On appelle polynôme irréductible tout polynôme de degré supérieur ou égal à 1, divisible uniquement par 1 ou par lui-même (à une constante multiplicative près). Les polynômes irréductibles jouent dans $\mathbb{K}[X]$ le même rôle que les nombres premiers dans $\mathbb{Z}$.

Exemple:

Calculer le PGCD de $x^3 + 2x^2 - x - 2$ et de $x^2 + 4x + 3$

On a :

$$\begin{aligned}x^3 + 2x^2 - x - 2 &= (x^2 + 4x + 3)(x - 2) + 4x + 4 \\x^2 + 4x + 3 &= (4x + 4)\left(\frac{x}{4} + \frac{3}{4}\right)\end{aligned}$$

Le PGCD est donc $4x + 4$ ou plutôt $x + 1$, l'habitude étant de donner le polynôme sous la forme normalisée ou unitaire.

Proposition 1.5.1 On a $(X^k)^{(n)} = A_k^{k-n}$ avec $A_k^n = n!C_n^k$

La démonstration se fait facilement par récurrence sur n .

1.6 Formules de Mac-Laurin et de Taylor

Théorème 1.6.1 (*Formule de Mac-Laurin*)

Soit $P \in \mathbb{K}_n[X]$, alors

$$P = \sum_{k=0}^n \frac{P^{(k)}(0)}{k!} X^k$$

Démonstration:

$$\text{Soit } P = \sum_{k=0}^n a_k X^k$$

$$\text{Pour tout entier } j, \quad P^{(j)} = \sum_{k=0}^n a_k (X^k)^{(j)} = \sum_{k=j}^n a_k A_k^j X^{k-j}$$

$$\text{Par suite, pour tout } j = 0, 1, \dots, n, \text{ on a } P^{(j)}(0) = a_j A_j^j = a_j j!$$

$$\text{et donc } a_j = \frac{P^{(j)}(0)}{j!}$$

$$\text{Par conséquent } P = \sum_{k=0}^n \frac{P^{(k)}(0)}{k!} X^k$$

Théorème 1.6.2 (*Formule de Taylor*)

Soit $P \in K_n[X]$ et $a \in \mathbb{K}$, alors $P = \sum_{k=0}^n \frac{P^{(k)}(a)}{k!} (X - a)^k$

Démonstration :

Il suffit d'appliquer la formule de Mac-Laurin au polynôme $P(X + a)$, on trouve:

$$P(X + a) = \sum_{k=0}^n \frac{P^{(k)}(a)}{k!} X^k$$

On obtient la formule souhaitée en substituant $X - a$ à X

1.7 Ordre de multiplicité d'une racine

Proposition 1.7.1 *Les propositions suivantes sont équivalentes*

- i) P est divisible par $(X - a)^k$ et pas par $(X - a)^{k+1}$
- ii) il existe Q tel que $Q(a) \neq 0$ et $P = (X - a)^k Q$
- iii) $P(a) = P'(a) = \dots = P^{(k-1)}(a) = 0$ et $P^{(k)}(a) \neq 0$

On dit que a est une racine de multiplicité k du polynôme P .

Démonstration:

i) $\implies$ ii)

Si P est divisible par $(X - a)^k$, il existe Q tel que $P = (X - a)^k Q$. Si on avait $Q(a) = 0$, alors Q pourrait se factoriser par $X - a$ et P serait divisible par $(X - a)^{k+1}$.

ii) $\implies$ iii)

Si $P = (X - a)^k Q$ avec $Q(a) \neq 0$, alors, on a, pour i compris entre 0 et k

$$P^{(i)}(X) = (X - a)^{k-i} Q^{(i)}(X) \text{ avec } Q(a) \neq 0$$

Ce résultat se montre aisément par récurrence. Il est vrai pour $i = 0$, et s'il est vrai pour $i < k$, alors :

$$\begin{aligned} P^{(i+1)}(X) &= (k - i)(X - a)^{k-i-1}(Q(X) + (X - a)Q')(X) \\ &= (X - a)^{k-i-1} Q_{i+1}(X) \end{aligned}$$

avec $Q_{i+1}(X) = (k - i)(Q(X) + (X - a)Q'(X))$

On a bien $P^{(i)}(a) = 0$ pour $0 \leq i \leq k-1$, et $P^{(k)}(a) = Q_k(a)$ différent de 0.

iii) $\implies$ i)

On applique la formule de Taylor et on factorise par $(X - a)^k$.

1.8 Théorème de d'Alembert

Théorème 1.8.1 (*admis*)

Tout polynôme non nul admet au moins une racine sur $\mathbb{C}$

Il en résulte que les polynômes irréductibles sont tous de degré 1, et que tout polynôme à coefficients complexes peut se factoriser sous la forme

$$\lambda \prod_{i \geq 0} (X - a_i)^{k_i}$$

Si $P = \sum_{i \geq 0} a_i X^i$, notons $\overline{P} = \sum_{i \geq 0} \overline{a_i} X^i$. Si z est complexe, on a alors : $\overline{P(z)} = \overline{P}(\overline{z})$ de sorte que si z est racine de P , alors $\overline{z}$ est racine de $\overline{P}$, avec le même ordre de multiplicité. Si P est à coefficients réels, alors $P = \overline{P}$, et si z est racine de P , alors $\overline{z}$ aussi. Les polynômes P à coefficients réels se décomposent alors sur $\mathbb{C}$ sous la forme :

$$P = \lambda \prod_{i \geq 0} (X - a_i)^{k_i} \prod_{i \geq 0} (X - z_i)^{m_i} (X - \overline{z_i})^{m_i}$$

et sur $\mathbb{R}$, en regroupant les parties conjuguées :

$$P = \lambda \prod_{i \geq 0} (X - a_i)^{k_i} \prod_{i \geq 0} (X^2 - \alpha_i X + \beta_i)^{m_i} \text{ avec } \alpha_i = 2\operatorname{Re}(z_i) \text{ et } \beta_i = |z_i|^2$$

Les polynômes irréductibles sur $\mathbb{R}$ sont donc de degré 1 ou 2.

Exemple : $X^4 + 1$ se factorise sur $\mathbb{R}$ sous la forme :

$$(X^2 - \sqrt{2}X + 1)(X^2 + \sqrt{2}X + 1)$$

1.9 Division suivant les puissances croissantes

Il s'agit d'une forme de division autre que la classique division euclidienne: si A et B sont deux polynômes, le terme constant de B étant non nul, et si n est un entier strictement positif, alors il existe des polynômes Q et R (déterminés de manière unique) tels que :

- $A = BQ + X^{n+1}R$.
- $\deg(Q)$ est inférieur ou égal à n .

On pose cette division un peu comme la classique division euclidienne, mais en écrivant les polynômes suivant les puissances croissantes, et en cherchant à éliminer d'abord les termes constants, puis les termes en X , etc...

Exemple: $A = 1+X$, $B = 1-X$, on trouve à l'ordre 2: $Q = 1+2X+2X^2$ et $R = 2$

$\begin{array}{r} 1+x \\ -1-x \\ \hline 2x \\ -2x-2x^2 \\ \hline 2x^2 \\ -2x^2-2x^3 \\ \hline 2x^3 \end{array}$	$\begin{array}{r} 1-x \\ \hline 1+2x+2x^2 \end{array}$
---	--

La division suivant les puissances croissantes est par exemple utilisée pour obtenir les décompositions en éléments simples des fractions rationnelles.

Chapitre 2

Fractions rationnelles

Définition 2.0.1 Une fraction rationnelle est le quotient de deux polynômes $\frac{A}{B}$ avec $B \neq 0$. On dit que deux fractions rationnelles $\frac{A}{B}$ et $\frac{C}{D}$ sont égales si et seulement si $AD = BC$ (comme dans $\mathbb{Q}$ pour les entiers). On dit que la fraction est irréductible si les deux polynômes A et B sont premiers entre eux. On note $\mathbb{k}(X)$ l'ensemble des fractions rationnelles de polynômes à coefficients dans $\mathbb{k}$, il n'est pas difficile de vérifier qu'il s'agit d'un corps.

2.1 Degré, partie entière

Définition 2.1.1 Soit $F = \frac{A}{B}$ une fraction rationnelle non nulle. L'entier relatif $\deg(F) = \deg(A) - \deg(B)$ est appelé degré de F .

Remarque

-Si $F = \frac{A}{B} = \frac{C}{D}$ (donc $AD = BC$), on a bien sur $\deg(A) - \deg(B) = \deg(C) - \deg(D)$

-Si F est en fait un polynôme, son degré (en tant qu'élément de $\mathbb{k}(X)$) coïncide avec son degré (en tant qu'élément de $\mathbb{k}[X]$).

-On étend la définition en posant que le degré de $F = 0$ est $-\infty$

-On vérifie : $\deg(F + G) \leq \max(\deg(F), \deg(G))$ et $\deg(FG) = \deg(F) + \deg(G)$.

Proposition 2.1.1 Tout élément F de $\mathbb{k}(X)$ s'écrit de manière unique $F = P + G$, où P est un polynôme, et où G est une fraction rationnelle de degré strictement négatif.

On dit que le polynôme P est la partie entière de la fraction rationnelle F .

Remarques et propriétés

- Posons $F = \frac{A}{B}$ et soit $A = BQ + C$ la division euclidienne de A par B .
On a l'égalité $F = Q + \frac{C}{B}$ avec $\deg(C) < \deg(B)$
La partie entière de F est donc le quotient dans la division de A par B .
En particulier, si $\deg(A) = \deg(B)$; la partie entière de F est la constante obtenue par quotient des coefficients dominants des polynômes A et B .
- Notons $E(F)$ la partie entière de toute fraction rationnelle F .
 - Pour tout polynôme P , on a $E(P) = P$.
 - Soit F dans $\mathbb{k}(X)$. On a $E(F) = 0 \iff$ le degré de F est strictement négatif..
 - Soit $F = \frac{A}{B}$, avec $\deg(A) \geq \deg(B)$. Alors $\deg(E(F)) = \deg(A) - \deg(B)$.

2.2 Pôles et partie polaires

Définition 2.2.1 (*pôle d'une fraction rationnelle*)

Soit $F = \frac{A}{B}$ une fraction rationnelle écrite sous forme irréductible
Soit α un élément de $\mathbb{k}$, et m dans $\mathbb{N}^*$. On dit que α est un pôle de F , avec la multiplicité m , si α est une racine du polynôme B avec la multiplicité m .

Remarques:

- Avec ces notations, α n'est pas racine de A car $A \wedge B = 1$.
- α est un pôle de F avec la multiplicité $m \iff F = \frac{A}{(X-\alpha)^m Q}$ avec $\begin{cases} A(\alpha) \neq 0 \\ Q(\alpha) \neq 0 \end{cases}$
- On parle de pôle simple si $m = 1$, et de pôle multiple si $m > 1$.
- On parle de pôle double si $m = 2$, triple si $m = 3$, etc..
- α est pôle simple de $F \iff F = \frac{A}{B}$ avec $A(\alpha) \neq 0, B(\alpha) = 0, B'(\alpha) \neq 0$.

De même α est un pôle double de $F \iff F = \frac{A}{B}$ avec $\begin{cases} A(\alpha) \neq 0 \\ B(\alpha) = B'(\alpha) = 0, B''(\alpha) \neq 0 \end{cases}$

- Soit F une fraction rationnelle à coefficients réels.
- Soit α un pôle complexe non réel de F , avec la multiplicité m .
- Alors $\bar{\alpha}$ est un pôle de F , avec la multiplicité m .

Proposition 2.2.1 (*Partie polaire*)

Soit F une fraction rationnelle de $\mathbb{k}(X)$, admettant α comme pôle de multiplicité $m \geq 1$.

Alors F s'écrit de manière unique $F = \sum_{k=1}^m \frac{\lambda_k}{(X-\alpha)^k} + G$ où $(\lambda_1, \dots, \lambda_m) \in \mathbb{k}^m$ et où G est une fraction rationnelle n'admettant pas α pour pôle. On dit alors que $\sum_{k=1}^m \frac{\lambda_k}{(X-\alpha)^k}$ est la partie polaire de F relativement au pôle α .

Remarque

Les pôles de G sont ceux de F (sauf α); avec les mêmes multiplicités respectives.

Cas d'un pôle simple

Soit F un élément de $\mathbb{k}(X)$ admettant α comme pôle simple. Il existe donc λ dans $\mathbb{k}$ tel que $F = \frac{\lambda}{X-\alpha} + G$, où α n'est pas un pôle de G . Voici deux méthodes permettant de calculer le coefficient λ .

- Posons $F = \frac{A}{(X-\alpha)Q}$, avec $A(\alpha) \neq 0$ et $Q(\alpha) \neq 0$. Alors $\lambda = \frac{A(\alpha)}{Q(\alpha)}$

Dans la pratique, on multiplie F par $(X - \alpha)$, et après simplification on substitue α à X .

Cette méthode est très adaptée au cas courant où B est factorisé.

- Posons $F = \frac{A}{B}$ avec $A(\alpha) \neq 0$ Alors $\lambda = \frac{A(\alpha)}{B'(\alpha)}$

Cette méthode est très adaptée au cas où B n'est pas factorisé.

Un cas classique est $B = X^n - 1$: les pôles sont les racines n -ièmes de l'unité.

Cas d'un pôle double

Soit F un élément de $\mathbb{k}(X)$ admettant α comme pôle double.

Il existe donc λ, μ dans $\mathbb{k}$ tels que $F = \frac{\lambda}{(X-\alpha)^2} + \frac{\mu}{X-\alpha} + G$, où α n'est pas un pôle de G .

- Posons $F = \frac{A}{(X-\alpha)^2 Q}$ avec $A(\alpha) \neq 0$ et $Q(\alpha) \neq 0$. Alors $\lambda = \frac{A(\alpha)}{Q(\alpha)}$
- Posons $F = \frac{A}{B}$, avec $A(\alpha) \neq 0$. Alors $\lambda = \frac{2A(\alpha)}{B''(\alpha)}$
- Une fois calculé le coefficient λ , on peut écrire : $H = F - \frac{\lambda}{(X-\alpha)^2} = \frac{\mu}{X-\alpha} + G$

Ou bien α n'est pas un pôle de H (donc $\mu = 0$ est c'est fini), ou bien α est un pôle simple de H et on est ramené à des méthodes connues.

Cas d'un pôle multiple

Soit F un élément de $\mathbb{K}(X)$ admettant α comme pôle avec la multiplicité $m \geq 2$.

Il existe $\lambda_1, \dots, \lambda_m$ dans $\mathbb{K}$ tel que $F = \sum_{k=1}^m \frac{\lambda_k}{(X-\alpha)^k} + G$ ou α n'est pas un pôle de G

On peut assez facilement calculer le coefficient λ_m de $\frac{1}{(X-\alpha)^m}$:

- Posons $F = \frac{A}{(X-\alpha)^m Q}$, avec $A(\alpha) \neq 0$ et $Q(\alpha) \neq 0$ alors $\lambda_m = \frac{A(\alpha)}{Q(\alpha)}$
- Une fois calculé λ_m , on peut écrire : $H = F - \frac{\lambda_m}{(X-\alpha)^m} = \sum_{k=1}^{m-1} \frac{\lambda_k}{(X-\alpha)^k} + G$
- On est alors en mesure de calculer $\lambda_{m-1}, \lambda_{m-2}, \dots$ etc. Cette méthode est cependant très lourde si m est "grand" (ce qui heureusement est rare)

Pour de "grandes" valeurs de m , on revient à $F = \frac{A}{(X-\alpha)^m Q}$ avec $\begin{cases} A(\alpha) \neq 0 \\ Q(\alpha) \neq 0 \end{cases}$

L'égalité $F = \sum_{k=1}^m \frac{\lambda_k}{(X-\alpha)^k} + G$ devient : $\frac{A}{Q} = \sum_{k=1}^m \lambda_k (X-\alpha)^{m-k} + (X-\alpha)^m G$

La substitution qui consiste à remplacer X par $X + \alpha$ donne alors :

$$A(\alpha + X) = (\lambda_m + \lambda_{m-1}X + \dots + \lambda_1 X^{m-1})Q(\alpha + X) + X^m G(X + \alpha)$$

On peut alors calculer successivement $\lambda_m, \dots, \lambda_2, \lambda_1$, par une méthode de division suivant les puissances croissantes du polynôme $A(\alpha + X)$ par le polynôme $Q(\alpha + X)$.

2.3 Décomposition en éléments simples

Proposition 2.3.1 (décomposition dans $\mathbb{C}(X)$)

Soient $F = \frac{A}{B}$ une fraction rationnelle à coefficients complexes.

Soit les $\alpha_1, \dots, \alpha_p$ pôles distincts de F , avec les multiplicités $r_1, \dots, r_p$

Alors F s'écrit de manière unique $F = E + \sum_{k=1}^p \left(\sum_{j=1}^{r_k} \frac{\lambda_{k,j}}{(X-\alpha_k)^j} \right)$

où E est la partie entière de F et où les $\lambda_{k,j}$ sont des éléments de $\mathbb{C}$.

Cette écriture est appelée décomposition en éléments simples de F dans $\mathbb{C}(X)$.

Remarque

Dans cette écriture, chaque somme $\sum_{j=1}^{r_k} \frac{\lambda_{k,j}}{(X-\alpha_k)^j}$ est la partie polaire de F pour α_k .

Proposition 2.3.2 (décomposition dans $\mathbb{R}(X)$)

Soit $F = \frac{A}{B}$ une fraction rationnelle de coefficients réels, sous forme irréductible.

Soit $B = \lambda \prod_{k=1}^p (X - \alpha_k)^{r_k} \prod_{k=1}^q (X^2 + b_k X + c_k)^{s_k}$ la factorisation de B dans $\mathbb{R}[X]$

Alors la fraction rationnelle F s'écrit de manière unique :

$$F = E + \sum_{k=1}^p \left(\sum_{j=1}^{r_k} \frac{\lambda_{k,j}}{(X-\alpha_k)^j} \right) + \sum_{k=1}^q \left(\sum_{j=1}^{s_k} \frac{c_{k,j} X + d_{k,j}}{(X^2 + b_k X + c_k)^j} \right)$$

où E est la partie entière de F et où les $\lambda_{k,j}, c_{k,j}, d_{k,j}$ sont des éléments de $\mathbb{R}$.

Cette écriture est appelée décomposition en éléments simples de F dans $\mathbb{R}(X)$

Remarques

- Dans la décomposition de la fraction rationnelle F dans $\mathbb{R}(X)$.
 - Les fractions $\frac{\lambda_{k,j}}{(X-\alpha_k)^j}$ sont appelées éléments simples de première espèce.
 - Les fractions $\frac{c_{k,j} X + d_{k,j}}{(X^2 + b_k X + c_k)^j}$ sont appelées éléments simples de seconde espèce.
- Ce qui a été dit pour les parties polaires permet:
 - D'effectuer complètement une décomposition dans $\mathbb{C}(X)$.
 - De calculer les éléments simples de première espèce dans $\mathbb{R}(X)$.

- Soit $F = \frac{A}{(X^2+bX+c)^m Q} \in \mathbb{R}(X)$, avec $b^2 - 4c < 0$, (A, Q) non divisibles par $X^2 + bX + c$

Alors $F = G + \sum_{j=1}^m \frac{c_j X + d_j}{(X^2+bX+c)^j}$, en groupant dans G ce qui ne relève pas de $(X^2 + bX + c)$.

Les coefficients les plus facilement "accessibles" sont c_m, d_m .

On peut en effet écrire $A = (c_m X + d_m)Q + (X^2 + bX + c)R(\dots)$

Soit ω une des deux racines complexes non réelles de $X^2 + bX + c$.

Si on substitue ω à X , on trouve: $A(\omega) = (c_m \omega + d_m)Q(\omega)$

On peut alors trouver c_m et d_m par identification:

- En utilisant la valeur de ω si elle est simple, notamment si $\omega = i$ ou $\omega = j$.
- Par abaissement successifs du degré en utilisant $\omega^2 = -b\omega - c$.

Une fois connus c_m et d_m , on considère $F - \frac{c_m X + d_m}{(X^2+bX+c)^m}$ et on cherche c_{m-1}, d_{m-1} , etc .

2.4 Pratique de la décomposition en éléments simples

Les méthodes vues précédemment permettent en principe de former les décompositions en éléments simples dans $\mathbb{R}(X)$ ou dans $\mathbb{C}(X)$.

Néanmoins, un certain nombre de techniques facilitent le travail:

- En diminuant globalement le nombre de coefficients à calculer.
- En permettant de calculer des coefficients peu facilement "accessibles", à condition cependant qu'il ne reste à ce stade que peu d'inconnues

Décomposition dans $\mathbb{C}(X)$ d'une fraction à coefficients réels

Soit $F = \frac{A}{B}$ un élément de $\mathbb{R}(X)$.

On peut considérer F comme un élément de $\mathbb{C}(X)$ et la décomposer en tant que telle.

Tout comme F , cette décomposition doit être invariante par conjugaison.

Il en résulte par exemple que la partie entière de F est un polynôme réel.

Il en résulte également que les parties polaires sont conjuguées deux à deux. Plus précisément, si α et $\bar{\alpha}$ sont deux pôles conjugués non réels de F , de multiplicité m , les parties polaires s'écrivent respectivement: $\sum_{k=1}^m \frac{\lambda_k}{(X-\alpha)^k}$ et

$$\sum_{k=1}^m \frac{\bar{\lambda}_k}{(X-\bar{\alpha})^k}$$

Cette idée permet donc de diminuer de moitié environ le nombre d'inconnues.

Utilisation de la parité ou de l'imparité

Si une fraction rationnelle est paire ou impaire, sa décomposition doit refléter cette propriété. On exprime cette invariance par les transformations $X \longrightarrow F(-X)$ ou $X \longrightarrow -F(-X)$, et on en déduit des relations sur les coefficients (le nombre d'inconnues diminue environ de moitié).

Injection de valeurs particulières

Quand il reste peu de coefficients à calculer, il peut être intéressant d'injecter, dans l'égalité entre F et sa décomposition, une ou plusieurs valeurs qui ne soient pas des pôles de F .

Si F est dans $\mathbb{R}(X)$, on peut injecter une valeur complexe comme i (ou j), l'identification donnant alors deux relations entre les coefficients réels inconnus.

Utilisation de la méthode $\lim_{x \rightarrow \infty} xF(x)$

On suppose ici que le degré de F est strictement négatif (la partie entière est donc nulle).

La décomposition de F fait apparaître des termes du type $\frac{\lambda_k}{X-\alpha_k}$ ou $\frac{a_k X + b_k}{X^2 + \beta_k X + \gamma_k}$. Le calcul de $\lim_{x \rightarrow \infty} xF(x)$ donne alors une relation liant les coefficients λ_k et a_k .

Cette méthode est intéressante quand il ne reste que un ou deux coefficients à calculer.

Exemples de référence

Décomposer $F = \frac{1}{X^n - 1}$ dans $\mathbb{C}(X)$

Ici $F = \frac{A}{B}$, avec $\begin{cases} A = 1 \\ B = X^n - 1 \end{cases}$ Les pôles (simples) sont les racines n -ièmes ω_k de 1

Pour tout k de $\{0, \dots, n-1\}$, on a $\frac{A(\omega_k)}{B'(\omega_k)} = \frac{1}{n\omega_k^{n-1}} = \frac{\omega_k}{n}$

La décomposition en éléments simples de F s'écrit donc

$$\frac{1}{X^n - 1} = \frac{1}{n} \sum_{k=0}^{n-1} \frac{\omega_k}{X - \omega_k}$$

Décomposer $F = \frac{1}{X(X+1)(X+2)\dots(X+n)}$ dans $\mathbb{R}(X)$.

Les pôles sont $0, -1, -2, \dots, -n$. Ce sont des pôles simples

La forme de la décomposition est : $F = \sum_{k=0}^n \frac{\lambda_k}{X+k}$

λ_k s'obtient en multipliant F par $X+k$ et en substituant $-k$ à X . On trouve :

$$\lambda_k = \prod_{j \neq k} \frac{1}{-k+j} = \prod_{j=0}^{k-1} \frac{1}{j-k} \prod_{j=k+1}^n \frac{1}{j-k} = (-1)^k \prod_{i=1}^k \frac{1}{i} \prod_{i=1}^{n-k} \frac{1}{i} = \frac{(-1)^k}{k!(n-k)!} = \frac{(-1)^k}{n!} C_n^k$$

Conclusion : $F = \frac{1}{X(X+1)(X+2)\dots(X+n)} = \frac{1}{n!} \sum_{k=0}^n \frac{(-1)^k C_n^k}{X+k}$

Décomposer $F = \frac{1}{X^3(X^2-1)}$ dans $\mathbb{R}(X)$

Les pôles sont 0 (triple), 1 (simple) et -1 (simple). La partie entière est nulle.

Posons. $F = \frac{a}{X^3} + \frac{b}{X^2} + \frac{c}{X} + \frac{d}{X-1} + \frac{e}{X+1}$. L'imparité de F donne $b = 0$ et $e = d$

On trouve a en multipliant par X^3 et en substituant 0 à X . Donc $a = -1$

On trouve d en multipliant par $(X-1)$ et en substituant 1 à X . Donc $d = \frac{1}{2}$

Si on utilise la méthode $\lim_{x \rightarrow \infty} xF(x)$, on trouve $0 = c + 1$. Donc $c = -1$

Finalement : $F = \frac{1}{X^3(X^2-1)} = -\frac{1}{X^3} - \frac{1}{X} + \frac{1}{2(X-1)} + \frac{1}{2(X+1)}$

Décomposer $F = \frac{1}{X(X^2+1)(X^2+X+1)(X^2-X+1)}$ dans $\mathbb{R}(X)$

La décomposition est de la forme: $F = \frac{a}{X} + \frac{cX+d}{X^2+1} + \frac{\alpha X+\beta}{X^2+X+1} + \frac{\lambda X+\mu}{X^2-X+1}$

L'imparité de F donne immédiatement : $d = 0, \mu = -\beta$ et $\lambda = \alpha$.

On cherche donc a, c, α, β tels que : $F = \frac{a}{X} + \frac{cX}{X^2+1} + \frac{\alpha X+\beta}{X^2+X+1} + \frac{\alpha X-\beta}{X^2-X+1}$

On trouve a en multipliant par X et en substituant 0 à X . Donc $a = 1$.

On trouve c en multipliant par X^2+1 et en substituant i à X . Donc $c = -1$.

On trouve α, β en multipliant par X^2+X+1 et en substituant j à X :

$$(X^2+X+1)F = \frac{1}{X(X^2+1)(X^2-X+1)} = \alpha X + \beta + (X^2+X+1)(\dots)$$

$$\implies \alpha j + \beta = \frac{1}{j(j^2+1)(j^2-j+1)} = \frac{1}{2} \implies \alpha = 0, \beta = \frac{1}{2}$$

Finalemment : $F = \frac{1}{X} - \frac{X}{X^2+1} + \frac{1}{2(X^2+X+1)} - \frac{1}{2(X^2-X+1)}$

Décomposer $F = \frac{X^5}{(X-1)^4}$ dans $\mathbb{R}(X)$

On écrit $X^5 = (X - 1 + 1)^5 = (X - 1)^5 + 5(X - 1)^4 + 10(X - 1)^3 + 10(X - 1)^2 + 5(X - 1) + 1$

On en déduit: $F = X + 4 + \frac{10}{(X-1)} + \frac{10}{(X-1)^2} + \frac{5}{(X-1)^3} + \frac{1}{(X-1)^4}$

Décomposer $F = \frac{X^8}{(X^2-X+1)^3}$ dans $\mathbb{R}(X)$

On procède à des divisions successives par $B = X^2 - X + 1$

$X^8 = Q_1B + R_1$, avec $Q_1 = X^6 + X^5 - X^3 - X^2 + 1$ et $R_1 = X - 1$.

$Q_1 = Q_2B + R_2$, avec $Q_2 = X^4 + 2X^3 + X^2 - 2X - 4$ et $R_2 = -2X + 5$

Ainsi: $X^8 = R_1 + R_2B + R_3B^2 + Q_3B^3$ puis $F = \frac{X^8}{B^3} = Q_3 + \frac{R_3}{B} + \frac{R_2}{B^2} + \frac{R_1}{B^3}$

Conclusion : $F = X^2 + 3X + 3 - \frac{2X+7}{X^2-X+1} - \frac{2X-5}{(X^2-X+1)^2} + \frac{X-1}{(X^2-X+1)^3}$